

A busca científica para entender,
aprimorar e potencializar a mente

MICHIO KAKU



O FUTURO DA MENTE

ROCCO HAN

VISIT...

LANZAROTE
Caliente.COM

DADOS DE COPYRIGHT

Sobre a obra:

A presente obra é disponibilizada pela equipe [Le Livros](#) e seus diversos parceiros, com o objetivo de oferecer conteúdo para uso parcial em pesquisas e estudos acadêmicos, bem como o simples teste da qualidade da obra, com o fim exclusivo de compra futura.

É expressamente proibida e totalmente repudiável a venda, aluguel, ou quaisquer uso comercial do presente conteúdo

Sobre nós:

O [Le Livros](#) e seus parceiros disponibilizam conteúdo de domínio público e propriedade intelectual de forma totalmente gratuita, por acreditar que o conhecimento e a educação devem ser acessíveis e livres a toda e qualquer pessoa. Você pode encontrar mais obras em nosso site: [LeLivros.site](#) ou em qualquer um dos sites parceiros apresentados [neste link](#)

"Quando o mundo estiver unido na busca do conhecimento, e não mais lutando por dinheiro e poder, então nossa sociedade poderá enfim evoluir a um novo nível."



MICHIO KAKU

**O FUTURO
DA MENTE**

A BUSCA CIENTÍFICA PARA ENTENDER,
APRIMORAR E POTENCIALIZAR A MENTE

Tradução de Angela Lobo

ROCCO ITALIA

*Este livro é dedicado à minha amada mulher, Shizue, e às minhas filhas,
Michelle e Alyson*

SUMÁRIO

Para pular o Sumário, clique [aqui](#).

AGRADECIMENTOS

INTRODUÇÃO

LIVRO I - A MENTE E A CONSCIÊNCIA

[1 | Desvendando a mente](#)

[2 | Consciência – O ponto de vista de um físico](#)

LIVRO II - A MENTE ACIMA DA MATÉRIA

[3 | Telepatia – Um doce pelo que você está pensando](#)

[4 | Telecinesia – A mente controlando a matéria](#)

[5 | Memórias e pensamentos feitos sob medida](#)

[6 | O cérebro de Einstein e a expansão de nossa inteligência](#)

LIVRO III - ALTERAÇÕES DE CONSCIÊNCIA

[7 | Em seus sonhos](#)

[8 | A mente pode ser controlada?](#)

[9 | Estados alterados de consciência](#)

[10 | A mente artificial e a consciência de silício](#)

[11 | A engenharia reversa do cérebro](#)

[12 | O futuro: a mente além da matéria](#)

[13 | A mente como energia pura](#)

[14 | A mente alienígena](#)

[15 | Observações finais](#)

APÊNDICE

NOTAS

LEITURA SUGERIDA

CRÉDITOS DAS ILUSTRAÇÕES

CRÉDITOS

O AUTOR

AGRADECIMENTOS

Foi um imenso prazer fazer entrevistas e interagir com esses notáveis cientistas, todos líderes em seus campos de pesquisa. Gostaria de agradecer-lhes a gentileza de terem me cedido seu tempo para entrevistas e discussões sobre o futuro da ciência. Além de contribuírem com sólidos fundamentos de suas respectivas áreas, deram-me orientações e inspiração.

Gostaria de agradecer a esses pioneiros e desbravadores, particularmente aqueles que concordaram em participar de meus programas especiais para os canais de televisão BBC, Discovery e Discovery Science, bem como de meus programas de rádio, *Science Fantastic* e *Explorations*.

Peter Doherty, ganhador do Prêmio Nobel, St. Jude Children's Research Hospital

Gerald Edelman, ganhador do Prêmio Nobel, Scripps Research Institute

Leon Lederman, ganhador do Prêmio Nobel, Illinois Institute of Technology

Murray Gell-Mann, ganhador do Prêmio Nobel, Santa Fe Institute e Cal Tech

Henry Kendall, ganhador do Prêmio Nobel, MIT, *in memoriam*

Walter Gilbert, ganhador do Prêmio Nobel, Universidade de Harvard

David Gross, ganhador do Prêmio Nobel, Kavli Institute for Theoretical Physics

Joseph Rotblat, ganhador do Prêmio Nobel, St. Bartholomew's Hospital

Yoichiro Nambu, ganhador do Prêmio Nobel, Universidade de Chicago

Steven Weinberg, ganhador do Prêmio Nobel, Universidade do Texas em Austin

Frank Wilczek, ganhador do Prêmio Nobel, MIT

Amir Aczel, autor de *Uranium Wars*

Buzz Aldrin, astronauta da Nasa, segundo homem a andar na Lua

Geoff Andersen, U.S. Air Force Academy, autor de *The Telescope*

Jay Barbree, autor de *Moon Shot*

John Barrow, físico, Universidade de Cambridge, autor de *Impossibility*

Marcia Bartusiak, autora de *Einstein's Unfinished Symphony*

Jim Bell, astrônomo na Cornell University

Jeffrey Bennet, autor de *Beyond UFOs*

Bob Berman, astrônomo, author *The Secrets of the Night Sky*

Leslie Biesecker, National Institutes of Health

Piers Bizony, autor de *How to Build Your Own Starship*

Michael Blaese, National Institutes of Health

Alex Boese, fundador do Museum of Hoaxes

Nick Bostrom, transumanista, Universidade de Oxford

Tenente-Coronel Robert Bowman, Institute for Space and Security Studies

Cynthia Breazeal, inteligência artificial, Laboratório de Mídia do MIT

Lawrence Brody, National Institutes of Health

Rodney Brooks, diretor do Laboratório de IA do MIT

Lester Brown, Earth Policy Institute

Michael Brown, astrônomo, Cal Tech

James Canton, autor de *The Extreme Future*

Arthur Caplan, diretor do Center for Bioethics da Universidade da Pensilvânia

Fritjof Capra, autor de *The Science of Leonardo*

Sean Carroll, cosmólogo, Cal Tech

Andrew Chaikin, autor de *A Man on the Moon*

Leroy Chiao, astronauta da Nasa

Eric Chivian, da organização International Physicians for the Prevention of Nuclear War

Deepak Chopra, autor de *Supercérebro*

George Church, diretor do Harvard's Center for Computational Genetics

Thomas Cochran, físico, Natural Resources Defense Council

Christopher Cokinos, astrônomo, autor de *Fallen Sky*

Francis Collins, National Institutes of Health

Vicki Colvin, nanotecnóloga, Universidade do Texas

Neal Comins, autor de *Hazards of Space Travel*

Steve Cook, porta-voz da Nasa

Christine Cosgrove, autora de *Normal at Any Cost*
Steve Cousins, CEO do Willow Garage Personal Robots Program
Phillip Coyle, ex-subsecretário de Defesa dos EUA
Daniel Crevier, AI, CEO da Coreco
Ken Crowell, astrônomo, autor de *Magnificent Universe*
Steven Cummer, ciência da computação, Duke University
Mark Cutkowsky, engenheiro mecânico, Universidade de Stanford
Paul Davies, físico, autor de *Superforce*
Daniel Dennet, filósofo, Tufts University
Michael Dertouzos, ciência da computação, MIT, *in memoriam*
Jared Diamond, ganhador do Prêmio Pulitzer, UCLA
Marriot DiChristina, revista *Scientific American*
Peter Dilworth, Laboratório de IA do MIT
John Donoghue, criador do Braingate, Universidade de Brown
Ann Druyan, viúva de Carl Sagan, Cosmos Studios
Freeman Dyson, Institute for Advanced Study, Universidade de Princeton
David Eagleman, neurocientista, Baylor College of Medicine
John Ellis, físico do CERN [Organização Europeia para a Pesquisa Nuclear]
Paul Erlich, ambientalista, Universidade de Stanford
Daniel Fairbanks, autor de *Relics of Eden*
Timothy Ferris, Universidade da Califórnia, autor de *Coming of Age in the Milky Way Galaxy*
Maria Finitzo, especialista em células-tronco, vencedora do Peabody Award
Robert Finkelstein, especialista em IA
Christopher Flavin, World Watch Institute
Louis Friedman, cofundador da Planetary Society
Jack Gallant, neurocientista, Universidade da Califórnia em Berkeley
James Garwin, cientista-chefe da Nasa
Evelyn Gates, autora de *Einstein's Telescope*
Michael Gazzaniga, neurologista, Universidade da Califórnia em Santa

Barbara

Jack Geiger, cofundador, Physicians for Social Responsibility

David Gelertner, cientista da computação, Universidade de Yale,
Universidade da Califórnia

Neal Gershenfeld, Laboratório de Mídia do MIT

Daniel Gilbert, psicólogo, Universidade de Harvard

Paul Gilster, autor de *Centauri Dreams*

Rebecca Goldberg, Environmental Defense Fund

Don Goldsmith, astrônomo, autor de *Runaway Universe*

David Goodstein, reitor adjunto da Cal Tech

J. Richard Gott III, Universidade de Princeton, autor de *Time Travel in
Einstein's Universe*

Stephen Jay Gould, biólogo, Universidade de Harvard, *in memoriam*

Embaixador Thomas Graham, satélites espiões e inteligência

John Grant, autor de *Corrupted Science*

Eric Green, National Institutes of Health

Ronald Green, autor de *Babies by Design*

Brian Greene, Universidade de Columbia, autor de *The Elegant Universe*

Alan Guth, físico, MIT, autor de *The Inflationary Universe*

William Hanson, autor de *The Edge of Medicine*

Leonard Hayflick, Escola de Medicina da Universidade da Califórnia

Donald Hillebrand, Argonne National Labs, o futuro do carro

Frank N. von Hippel, físico, Universidade de Princeton

Allan Hobson, psiquiatra, Universidade de Harvard

Jeffrey Hoffman, astronauta da NASA, MIT

Douglas Hofstadter, ganhador do Prêmio Pulitzer, Universidade de Indiana,
autor de *Gödel, Escher, Bach*

John Horgan, Stevens Institute of Technology, autor de *The End of Science*

Jamie Hyneman, apresentador da série *Os caçadores de mitos* [Discovery]

Chris Impey, astrônomo, autor de *The Living Cosmos*

Robert Irie, Laboratório de IA do MIT

P. J. Jacobowitz, revista *PC*

Jay Jaroslav, Laboratório de IA do MIT

Donald Johanson, antropólogo, descobridor do fóssil Lucy

George Johnson, jornalista de ciências no *The New York Times*

Tom Jones, astronauta da Nasa

Steve Kates, astrônomo

Jack Kessler, especialista em células-tronco, vencedor do Peabody Award

Robert Kirshner, astrônomo, Universidade de Harvard

Kris Koenig, astrônomo

Lawrence Krauss, Universidade do Estado do Arizona, autor de *Physics of Star Trek*

Lawrence Kuhn, cineasta e filósofo, *Closer to Truth*

Ray Kurzweil, inventor, autor de *The Age of Spiritual Machines*

Robert Lanza, biotecnologia, Advanced Cell Technologies

Roger Launius, autor de *Robots in Space*

Stan Lee, criador da Marvel Comics e do Homem-Aranha

Michael Lemonick, editor chefe de ciências da *Time*

Arthur Lerner-Lam, geólogo, vulcanista

Simon LeVay, autor de *When Science Goes Wrong*

John Lewis, astrônomo, Universidade do Arizona

Alan Lightman, MIT, autor de *Einstein's Dreams*

George Linehan, autor de *Space One*

Seth Lloyd, MIT, autor de *Programming the Universe*

Werner R. Loewenstein, ex-diretor do Cell Physics Laboratory, Universidade de Columbia

Joseph Lykken, físico, Fermi National Laboratory

Pattie Maes, Laboratório de IA do MIT

Robert Mann, autor de *Forensic Detective*

Michael Paul Mason, autor de *Head Cases: Stories of Brain Injury and Its Aftermath*

Patrick McCray, autor de *Keep Watching the Skies*
Glenn McGee, autor de *The Perfect Baby*
James McLurkin, Laboratório de IA do MIT
Paul McMillan, diretor do Space Watch
Fulvia Melia, astrônoma, Universidade do Arizona
William Meller, autor de *Evolution Rx*
Paul Meltzer, National Institutes of Health
Marvin Minsky, MIT, autor de *The Society of Minds*
Hans Moravec, autor de *Robot*
Phillip Morrison, físico, MIT, *in memoriam*
Richard Muller, astrofísico, Universidade da Califórnia em Berkeley
David Nahamoo, IBM Human Language Technology
Christina Neal, vulcanista
Miguel Nicolelis, neurocientista, Duke University
Shinji Nishimoto, neurologista, Universidade da Califórnia em Berkeley
Michael Novacek, American Museum of Natural History
Michael Oppenheimer, ambientalista, Universidade de Princeton
Dean Ornish, especialista em câncer e doenças do coração
Peter Palese, virologista, Mount Sinai School of Medicine
Charles Pellerin, oficial da Nasa
Sidney Perkowitz, autor de *Hollywood Science*
John Pike, GlobalSecurity.org
Jena Pincott, autora de *Os homens preferem mesmo as loiras?*
Steven Pinker, psicólogo, Universidade de Harvard
Thomas Poggio, MIT, Inteligência Artificial
Correy Powell, editor da revista *Discover*
John Powell, fundador da JP Aerospace
Richard Preston, autor de *Hot Zone* e *Demon in the Freezer*
Raman Prinja, astrônomo, University College London

David Quammen, biologia evolutiva, autor de *The Reluctant Mr. Darwin*

Katherine Ramsland, cientista forense

Lisa Randall, Universidade de Harvard, autora de *Warped Passages*

Sir Martin Rees, Astrônomo Real da Grã-Bretanha, Universidade de Cambridge, autor de *Before the Beginning*

Jeremy Rifkin, Foundation for Economic Trends

David Riquier, Laboratório de Mídia do MIT

Jane Rissler, Union of Concerned Scientists

Steven Rosenberg, National Institutes of Health

Oliver Sacks, neurologista, Universidade de Columbia

Paul Saffo, futurista, Institute of the Future

Carl Sagan, Cornell University, autor de *Cosmos, in memoriam*

Nick Sagan, coautor de *You Call This the Future?*

Michael H. Salamon, programa da Nasa Beyond Einstein

Adam Savage, apresentador da série *Os caçadores de mitos*

Peter Schwartz, futurista, fundador da Global Business Network

Michael Shermer, fundador da Skeptic Society e da revista *Skeptic*

Donna Shirley, programa Marte da Nasa

Seth Shostak, Instituto Seti

Neil Shubin, autor de *Your Inner Fish*

Paul Shurch, SETI League

Peter Singer, autor de *Wired for War*

Simon Singh, autor de *The Big Bang*

Gary Small, autor de *iBrain*

Paul Spudis, autor de *Odyssey Moon Limited*

Stephen Squyres, astrônomo, Cornell University

Paul Steinhardt, Universidade de Princeton, autor de *Endless Universe*

Jack Stern, cirurgia com células-tronco

Gregory Stock, UCLA, autor de *Redesigning Humans*

Richard Stone, autor de *NEOs e Tunguska*

Brian Sullivan, Hayden Planetarium

Leonard Susskind, físico, Universidade de Stanford

Daniel Tammet, autor de *Nascido em um dia azul*

Geoffrey Taylor, físico, Universidade de Melbourne

Ted Taylor, projetista de ogivas nucleares dos Estados Unidos, *in memoriam*

Max Tegmark, cosmólogo, MIT

Alvin Toffler, autor de *The Third Wave*

Patrick Tucker, World Future Society

Chris Turney, University of Wollongong, autor de *Ice, Mud and Blood*

Neil de Grasse Tyson, diretor do Hayden Planetarium

Sesh Velamoor, Foundation for the Future

Robert Wallace, autor de *Spycraft*

Kevin Warwick, ciborgues humanos, Universidade de Reading, Reino Unido

Fred Watson, astrônomo, autor de *Stargazer*

Mark Weiser, Xerox PARC, *in memoriam*

Alan Weisman, autor de *O mundo sem nós*

Daniel Wertheimer, SETI at Home, Universidade da Califórnia em Berkeley

Mike Wessler, Laboratório de IA do MIT

Roger Wiens, astrônomo, Los Alamos National Laboratory

Arthur Wiggins, autor de *The Joy of Physics*

Anthony Wynshaw-Boris, National Institutes of Health

Carl Zimmer, biólogo, autor de *Evolution*

Robert Zimmerman, autor de *Leaving Earth*

Robert Zubrin, fundador da Mars Society

Gostaria também de agradecer ao meu agente literário, Stuart Krichevsky, que esteve ao meu lado por todos esses anos, oferecendo bons conselhos sobre meus livros. O discernimento dele sempre me ajudou. Além disso, gostaria de agradecer a meus editores, Edward Kastenmeier e Melissa Danaczko, que me orientaram neste trabalho e forneceram conselhos editoriais inestimáveis. E também gostaria de agradecer à dra. Michelle Kaku, minha filha, neurologista residente no Mount Sinai Hospital em Nova York, pelas discussões estimulantes,

ponderadas e enriquecedoras sobre o futuro da neurologia. Sua leitura cuidadosa e aprofundada do manuscrito foi de grande valor para a forma e o conteúdo deste livro.

INTRODUÇÃO

Os dois maiores mistérios de toda a natureza são a mente e o universo. Com a nossa vasta tecnologia, fomos capazes de fotografar galáxias a bilhões de anos-luz de distância, manipular genes que controlam a vida e sondar os redutos mais escondidos do átomo. No entanto, a mente e o universo ainda nos ludibriam e intrigam. São os rincões mais misteriosos e fascinantes da ciência.

Para apreciar a grandiosidade do universo, basta olhar para o céu à noite, iluminado por bilhões de estrelas. Desde que nossos ancestrais se extasiaram pela primeira vez com o esplendor do céu estrelado, somos confrontados com as eternas perguntas: De onde veio tudo isso? O que tudo isso significa?

Para contemplar o mistério da mente, basta nos olharmos no espelho e tentar imaginar o que se esconde por trás de nossos olhos. Isso desperta perguntas assombrosas, como: Temos alma? O que acontece quando morremos? Quem sou “eu”, afinal? E o mais importante é que nos leva à pergunta derradeira: Qual é nossa parte no grande esquema cósmico? Como disse o famoso biólogo vitoriano Thomas Huxley, “A questão das questões para a humanidade, o problema existente por trás de todos os outros e também o mais interessante, é a definição do lugar do homem na Natureza e sua relação com o Cosmos”.

Há 100 bilhões de estrelas na Via Láctea, aproximadamente o mesmo número de neurônios que temos no cérebro. Seria preciso viajar 39 trilhões de quilômetros até a primeira estrela fora do sistema solar para encontrar algo tão complexo quanto o que temos em cima do pescoço. A mente e o universo constituem o maior de todos os desafios científicos, mas também apresentam uma relação curiosa. Por um lado, situam-se em polos opostos. Um concerne à vastidão do espaço sideral, onde encontramos estranhos habitantes, como buracos negros, explosões de estrelas e colisões de galáxias. O outro concerne ao espaço interior, onde estão nossas esperanças e desejos mais íntimos. A mente está logo ali no próximo pensamento, mas quando nos pedem para defini-la ou explicá-la, não temos a menor ideia.

Entretanto, apesar de opostos nesse aspecto, a mente e o universo também compartilham de uma mesma narrativa. Ambos estão envolvidos em superstição e magia desde tempos remotos. Astrólogos e frenólogos afirmaram encontrar o significado do universo em cada constelação do zodíaco e em cada protuberância do crânio. Adivinhos e videntes têm sido, ao mesmo tempo, enaltecidos e demonizados ao longo dos anos.

O universo e a mente continuam a se cruzar de várias maneiras, graças, em grande parte, a algumas ideias inovadoras que costumamos encontrar na ficção científica. Ao ler esses livros quando criança, eu sonhava em ser da Slan, uma raça de telepatas criada por A. E. van Vogt. Fascinava-me como um mutante chamado Mula era capaz de ativar seus vastos poderes telepáticos e quase

dominar o Império Galáctico na *Trilogia da Fundação*, de Isaac Asimov. E no filme *O planeta proibido*, eu tentava entender como uma civilização milhões de anos à frente da nossa conseguia canalizar seus enormes poderes telecinéticos para remodelar a realidade conforme seus anseios e desejos.

Quando eu tinha uns 10 anos, “O Incrível Dunninger” apareceu na televisão. O público ficava deslumbrado com seus espetaculares truques de mágica. Seu lema era: “Para quem acredita, nenhuma explicação é necessária; para quem não acredita, nenhuma explicação é suficiente.” Um dia ele declarou que iria enviar pensamentos a milhões de pessoas em todos os Estados Unidos. Fechou os olhos e se concentrou, dizendo que estava transmitindo o nome de um presidente do país. Pediu às pessoas que escrevessem num cartão-postal o nome que lhes viesse à mente e enviassem pelo correio. Na semana seguinte, anunciou, triunfante, que tinha recebido milhares de cartões com o nome “Roosevelt”, o mesmo que ele “transmitira” para todo o país.

Não fiquei impressionado. Na época, o legado de Roosevelt era muito forte para quem tinha atravessado a Depressão e a Segunda Guerra Mundial, portanto, não era surpresa. (Pensei que seria surpreendente se ele tivesse pensado no presidente Millard Fillmore.)

Ainda assim, aquilo alimentava minha imaginação e não resisti a fazer alguns experimentos de telepatia, concentrando-me o máximo possível para ler a mente das pessoas. Eu fechava os olhos, me concentrava intensamente tentando “ouvir” os pensamentos dos outros e mover por telecinesia os objetos do meu quarto.

Nunca consegui.

Talvez em algum lugar da Terra houvesse telepatas, mas eu não era um deles. No decorrer desse processo, comecei a perceber que provavelmente as incríveis façanhas dos telepatas eram impossíveis – pelo menos, sem ajuda externa. Mas, nos anos que se seguiram, também fui aprendendo outra lição: para penetrar nos maiores segredos do universo, não era necessário ter poderes telepáticos nem sobre-humanos. Era preciso apenas ter a mente aberta, determinada e curiosa. E sobretudo para saber se as fantásticas criações da ficção científica seriam possíveis, era preciso mergulhar na física avançada. Para identificar o ponto preciso onde o possível se torna o impossível, é necessário respeitar e entender as leis da física.

Duas paixões têm inflamado minha imaginação durante todos esses anos: entender as leis fundamentais da física e ver como a ciência vai configurar nossa vida no futuro. A fim de ilustrar isso e compartilhar meu entusiasmo em explorar as leis definitivas da física, escrevi os livros *Hiperespaço*, *Para além de Einstein* e *Mundos paralelos*. E para expressar meu fascínio pelo futuro, escrevi *Visões do futuro*, *Física do impossível* e *A física do futuro*. Durante a pesquisa e a escrita desses livros, fui constantemente lembrado de que a mente humana ainda é uma das forças mais poderosas e misteriosas desse mundo.

Na verdade, no curso da história, não fomos capazes de entender o que é a mente e como ela funciona. Os antigos egípcios, apesar de todas as grandes realizações nas artes e nas ciências, acreditavam que o cérebro era um órgão inútil e o jogavam fora quando embalsamavam os faraós. Aristóteles tinha certeza de que a alma residia no coração, e não no cérebro, cuja única função era resfriar o sistema cardiovascular. Outros, como Descartes, pensavam que a alma entrava no corpo através da minúscula glândula pineal, no cérebro. Mas, na falta de evidências sólidas, nenhuma dessas teorias pôde ser comprovada.

Essa “idade das trevas” persistiu por milhares de anos, e por bons motivos. O cérebro pesa apenas um quilo e meio, mas é o elemento mais complexo no sistema solar. Embora ocupe apenas 2% do peso do corpo, o cérebro tem um apetite voraz, consumindo 20% do total de nossa energia (em recém-nascidos, o cérebro consome 65% da energia), e 80% dos nossos genes são ativados no cérebro. A estimativa de neurônios existentes no crânio é de 100 bilhões, com quantidades exponenciais de conexões e caminhos neurais.

Em 1977, quando o astrônomo Carl Sagan escreveu o livro ganhador do Prêmio Pulitzer, *Os dragões do Éden*, ele resumiu o que se conhecia sobre o cérebro até aquele momento. O livro foi maravilhosamente bem escrito e buscou representar o que havia de mais desenvolvido na neurociência que, na época, se baseava fortemente em três fontes principais. A primeira era a comparação de nosso cérebro com os de outras espécies. Tal tarefa era enfadonha e difícil, porque envolvia a dissecação de cérebros de milhares de animais. A segunda metodologia era igualmente indireta, pois analisava vítimas de derrames e doenças, que passavam a apresentar um comportamento estranho. Somente após a morte, uma autópsia podia revelar qual parte do cérebro tinha alguma disfunção. Na terceira, os cientistas usavam eletrodos para investigar o cérebro e, lentamente e a duras penas, definiam qual parte do órgão influenciava determinado comportamento.

Mas os instrumentos básicos da neurociência não forneciam um modo sistemático de analisar o cérebro. Não se podia simplesmente recorrer a uma vítima de derrame com dano na área específica que se queria estudar. Dado que o cérebro é um sistema vivo, dinâmico, a autópsia quase nunca revelava os aspectos mais interessantes, como as formas de interação entre as diversas partes do órgão, ou como produziam pensamentos tão diversos quanto amor, ódio, ciúme e curiosidade.

REVOLUÇÕES GÊMEAS

Quatrocentos anos atrás foi inventado o telescópio e, praticamente de um dia

para o outro, esse novo instrumento miraculoso passou a vasculhar os corpos celestes. Trata-se de um dos instrumentos mais revolucionários (e perturbadores) de todos os tempos. De repente, podíamos ver com nossos próprios olhos os mitos e dogmas do passado evaporando-se como a bruma da manhã. Em vez de exemplos perfeitos da sabedoria divina, a Lua tinha crateras acidentadas, o Sol tinha manchas escuras, Júpiter tinha luas, Vênus tinha fases, Saturno tinha anéis. Aprendeu-se mais sobre o universo nos quinze anos seguintes à invenção do telescópio do que em toda a história da humanidade.

Assim como ocorreu com a invenção do telescópio, em meados dos anos 1990 e 2000 a introdução dos aparelhos de imagem por ressonância magnética (IRM) e de uma série de equipamentos avançados de varreduras cerebrais transformaram a neurociência. Aprendemos mais sobre o cérebro nos últimos quinze anos do que em toda a história humana anterior, e a mente, antes considerada fora de alcance, está finalmente assumindo o lugar central.

O ganhador do Prêmio Nobel Eric R. Kandel, do Max Planck Institut de Tübingen, Alemanha, diz: “Os *insights* mais valiosos sobre a mente humana que emergiram nesse período não vieram das disciplinas tradicionalmente relacionadas com a mente – filosofia, psicologia, psicanálise. Vieram de uma fusão dessas disciplinas com a biologia do cérebro...”

Os físicos tiveram um papel essencial nesse empreendimento, produzindo uma avalanche de instrumentos e siglas, como IRM, EEG, TEP, TAC, EMT, EET e ECP, que mudaram radicalmente o estudo do cérebro. Com essas máquinas, passamos a conseguir ver os pensamentos se movendo dentro do cérebro vivo e pensante. Como diz o neurologista V. S. Ramachandran, da Universidade da Califórnia, em San Diego: “Todas essas questões que os filósofos têm estudado há milênios, nós, cientistas, podemos começar a explorar fazendo imagens do cérebro, estudando pacientes, e fazendo as perguntas certas.”

Olhando para o passado, vejo que algumas das minhas incursões iniciais no mundo da física se cruzaram com essas mesmas tecnologias que hoje abrem a mente para a ciência. No ensino médio, por exemplo, tomei conhecimento de uma nova forma de matéria, chamada antimatéria, e decidi realizar um projeto de ciências sobre o tema. Como se trata de uma das substâncias mais incomuns existentes, precisei apelar para a antiga Comissão de Energia Atômica e consegui uma quantidade mínima de sódio-22, uma substância que emite naturalmente um elétron positivo (antielétron, ou pósitron). De posse da minha pequena amostra, consegui fazer uma câmara de nuvens e um poderoso campo magnético que me permitiram fotografar o rastro de vapor deixado pelas partículas de antimatéria. Eu não sabia na época, mas o sódio-22 logo iria se tornar essencial numa nova tecnologia chamada tomografia por emissão de pósitrons (TEP), que desde então tem nos revelado um conhecimento surpreendente sobre o cérebro pensante.

Outra tecnologia que experimentei na escola foi a ressonância magnética.

Assisti a uma palestra de Felix Bloch, da Universidade de Stanford, que recebeu o Prêmio Nobel de Física junto com Edward Purcell, em 1952, pela descoberta da ressonância magnética nuclear. O dr. Bloch ensinou a nós, jovens alunos, que dentro de um campo magnético forte os átomos se alinham verticalmente, como as agulhas da bússola. Se aplicamos a esses átomos um pulso de rádio numa frequência de ressonância específica, eles invertem o alinhamento. Quando retornam à posição inicial, eles emitem outro pulso, como um eco, que nos permite determinar a identidade desses átomos. Mais tarde, usei o princípio da ressonância magnética para construir um acelerador de partículas de 2,3 milhões de elétrons-volt na garagem da casa de minha mãe.

Poucos anos depois, quando eu era calouro da Universidade de Harvard, tive a honra de ser aluno de eletrodinâmica do dr. Purcell. Na mesma época, eu tinha um emprego de verão e surgiu a chance de trabalhar com o dr. Richard Ernst, que estava tentando generalizar o trabalho de Bloch e Purcell sobre ressonância magnética. Ele teve um sucesso espetacular, e acabou ganhando o Nobel de Física em 1991 por lançar as bases do aparelho moderno de IRM. Este, por sua vez, produziu fotografias detalhadas do cérebro vivo ainda mais precisas do que a varredura por TEP.

CAPACITANDO A MENTE

Eu me tornei professor de física teórica, mas minha fascinação pela mente permaneceu. Era sensacional ver que na última década os avanços da física possibilitaram algumas proezas do mentalismo que me extasiavam quando menino. Usando varredura por IRM, os cientistas conseguem ler os pensamentos que circulam em nosso cérebro. Podem também inserir um chip num cérebro totalmente paralisado e ligá-lo ao computador de modo que, apenas com o pensamento, o paciente consiga navegar na internet, ler e escrever e-mails, jogar videogames, controlar a cadeira de rodas, operar eletrodomésticos e manejar braços mecânicos. Por meio de um computador, esses pacientes podem fazer tudo o que faz uma pessoa normal.

Agora os cientistas estão indo ainda mais longe, conectando o cérebro diretamente a um exoesqueleto sobreposto aos membros paralisados do paciente. Um dia os tetraplégicos poderão ter uma vida muito próxima à normal. Esses exoesqueletos também podem nos dar superpoderes para lidar com emergências de vida ou morte. Talvez um dia os astronautas possam até explorar planetas, sentados confortavelmente na sala de casa, controlando mentalmente substitutos mecânicos.

Como no filme *Matrix*, talvez um dia possamos fazer o download de

lembranças e habilidades usando computadores. Em estudos com animais, os cientistas já conseguiram inserir memórias no cérebro dos bichos. Talvez seja apenas uma questão de tempo para que nós também sejamos capazes de inserir memórias em nosso cérebro e aprender coisas novas, passar as férias em novos lugares, adquirir novas habilidades. E se for possível transferir capacidades técnicas para a mente de trabalhadores e cientistas, a economia mundial pode ser beneficiada. Poderemos até compartilhar essas memórias. Talvez um dia os cientistas possam construir uma “internet da mente”, uma “mentenet”, que envie eletronicamente pensamentos e emoções para o mundo inteiro. Até os sonhos poderão ser gravados e enviados por “mentemail”, via internet.

A tecnologia também pode nos dar o poder de aumentar a inteligência. Já há progressos no entendimento dos poderes extraordinários dos *savants*, cujas capacidades mentais, artísticas e matemáticas são realmente impressionantes. Além disso, os genes que nos separam dos macacos estão sendo sequenciados, dando-nos uma ideia sem paralelo das origens evolucionárias do cérebro. Já foram isolados genes que podem aumentar a memória e o desempenho mental de animais.

O entusiasmo e o futuro promissor criado por esses avanços surpreendentes são de tal magnitude que já atraíram a atenção dos políticos. De fato, a ciência do cérebro de repente se tornou fonte de uma rivalidade transoceânica entre as maiores potências econômicas do planeta. Em janeiro de 2013, tanto o presidente norte-americano Barack Obama como a União Europeia anunciaram que poderiam destinar verbas de bilhões de dólares a dois projetos independentes para desenvolver a engenharia reversa do cérebro. Decifrar os intricados circuitos neurais no cérebro, antes considerados completamente fora do alcance da ciência moderna, é hoje o foco de dois projetos arrebatadores que, como o Projeto Genoma Humano, irão mudar o panorama médico e científico. Além de nos oferecer um conhecimento incomparável da mente humana, irá gerar novas indústrias, promover atividades econômicas e abrir novos horizontes para a neurociência.

Quando os caminhos neurais forem decodificados, poderemos finalmente compreender as origens exatas das doenças mentais, o que talvez leve à cura desse sofrimento tão antigo. Essa decodificação também tornará possível criar uma cópia do cérebro, o que levanta questões filosóficas e éticas. Quem somos nós, se nossa consciência pode ser transferida para um computador? Podemos brincar até com o conceito de imortalidade. Nosso corpo pode se deteriorar e morrer, mas nossa consciência viverá para sempre?

Além disso, talvez um dia, num futuro distante, a mente possa se libertar das restrições corporais e viajar pelas estrelas, como vários cientistas já especularam. Podemos imaginar que, daqui a séculos, talvez seja possível lançar no espaço raios laser com nossos diagramas neurais completos – talvez a

maneira mais conveniente de nossa consciência explorar as estrelas.

Surge hoje um cenário científico brilhante, que irá redefinir o destino humano. Estamos entrando na idade de ouro da neurociência.

Ao fazer tais previsões, tive a ajuda inestimável de cientistas que amavelmente me permitiram entrevistá-los, divulgar suas ideias em cadeia nacional de rádio, e até levar uma equipe de televisão a seus laboratórios. São esses cientistas que estão ditando as bases para o futuro da mente. Para incorporar suas ideias a este livro, fiz apenas duas exigências: (1) as previsões devem obedecer rigorosamente às leis da física, e (2) devem existir protótipos como prova de que as ideias funcionam.

ATINGIDO PELA DOENÇA MENTAL

Quando escrevi uma biografia de Einstein, intitulada *O cosmo de Einstein*, vi-me às voltas com detalhes minuciosos de sua vida privada. Eu sabia que o filho mais novo de Einstein sofria de esquizofrenia, mas não imaginava o enorme peso emocional que isso teve na vida do grande cientista. Einstein foi também afetado de outra maneira pela doença mental. Um de seus colegas mais chegados era o físico Paul Ehrenfest, que o ajudou a criar a teoria da relatividade geral. Após sofrer fases de depressão, Ehrenfest matou tragicamente o próprio filho, portador de síndrome de Down, e se suicidou. Com o passar dos anos, descobri que muitos colegas e amigos sofreram com doenças mentais na família.

A doença mental afetou profundamente a minha vida também. Muitos anos atrás, minha mãe morreu após uma longa batalha contra o mal de Alzheimer. Era de cortar o coração vê-la perder gradualmente a lembrança das pessoas amadas e olhá-la nos olhos e perceber que ela não sabia quem eu era. Eu via os resquícios de humanidade se extinguindo lentamente. Ela havia passado a vida inteira lutando para sustentar a família e, em vez de aproveitar a aposentadoria, todas as lembranças que valorizava lhe foram roubadas.

À medida que os *baby boomers* envelhecem, a triste experiência que eu e muitos outros tivemos se repetirá pelo mundo afora. Meu desejo é que o rápido avanço da neurociência possa um dia aliviar o sofrimento dos atingidos pela doença mental e pela demência.

O QUE ESTÁ CONDUZINDO ESSA REVOLUÇÃO?

Os dados provenientes das varreduras cerebrais estão sendo decodificados, e o

progresso é impressionante. Várias vezes por ano, vemos manchetes nos jornais anunciando uma nova descoberta. Da invenção do telescópio à era espacial passaram-se 350 anos, mas faz apenas quinze anos que a IRM e os mapeamentos cerebrais avançados foram introduzidos, permitindo conectarmos ativamente o cérebro ao mundo externo. *Por que tão rápido, e quanto ainda está por vir?*

Parte desse rápido progresso ocorreu porque os físicos de hoje têm um bom conhecimento do eletromagnetismo, que rege os sinais elétricos que percorrem os neurônios. As equações matemáticas de James Clerk Maxwell, usadas para calcular a física de antenas, radares, receptores de rádio e torres de micro-ondas, formam os alicerces da tecnologia de IRM. Levou séculos para se desvendar o segredo do eletromagnetismo, mas a neurociência saboreia os frutos dessa grande conquista.

No Livro I, faço um levantamento da história do cérebro e mostro como uma infinidade de novos instrumentos saiu dos laboratórios de física e nos deu imagens coloridas dos mecanismos do pensamento. Dado que a consciência tem um papel central em qualquer discussão sobre a mente, dou também a perspectiva de um físico, oferecendo uma definição de consciência que inclui o reino animal. Na verdade, trago uma classificação de consciências, mostrando que é possível atribuir números a vários tipos delas.

Mas para dar uma resposta completa à pergunta de como essa tecnologia vai avançar, é preciso retomar a lei de Moore, que afirma que a capacidade do computador duplica a cada dois anos. As pessoas se surpreendem diante do simples fato de que um telefone celular tem hoje mais capacidade de computação do que a Nasa *inteira* quando pôs dois homens na Lua em 1969. Hoje os computadores têm capacidade suficiente para registrar os sinais elétricos emitidos pelo cérebro e para decodificá-los parcialmente numa linguagem digital conhecida. Isso torna possível uma interface direta do cérebro com o computador para controlar qualquer objeto próximo. Esse campo, que cresce rapidamente, é chamado ICM (interface cérebro-máquina), e a tecnologia chave é o computador.

No Livro II exploro essa nova tecnologia, que possibilita o registro de memórias, leitura da mente, gravação de sonhos em vídeo, e telecinesia.

No Livro III, investigo formas alternativas de consciência, desde sonhos, drogas e doença mental a robôs e alienígenas do espaço sideral. Também falo sobre o potencial de controle e manejo do cérebro para lidar com doenças como depressão, mal de Parkinson, de Alzheimer e muitas outras. Discuto ainda o Projeto Brain, anunciado pelo presidente Obama, e o Projeto do Cérebro Humano, da União Europeia, que deverão destinar bilhões de dólares para decodificar os caminhos cerebrais até o nível neural. Esses dois programas certamente abrirão áreas de pesquisa inteiramente novas, trazendo novas maneiras de tratar a doença mental e revelando os segredos mais íntimos da

consciência.

Uma vez dada a definição de consciência, podemos usá-la para explorar também a consciência não humana (isto é, a consciência de robôs). Como serão os robôs avançados? Terão emoções? Poderão ser uma ameaça? E podemos explorar também a consciência de alienígenas, que podem ter metas totalmente diferentes das nossas.

No Apêndice, vou discutir a ideia talvez mais estranha de toda a ciência: o conceito da física quântica, de que a consciência pode ser a base fundamental da realidade.

Não faltam propostas para esse campo em expansão. Somente o tempo dirá quais delas são meras ilusões criadas pela imaginação fértil dos autores de ficção científica, e quais representam caminhos sólidos para a pesquisa científica. O progresso da neurociência tem sido astronômico e, em vários aspectos, a chave é a física moderna, que usa o poder total das forças eletromagnéticas e nucleares para investigar os grandes segredos ocultos em nossa mente.

É preciso enfatizar que não sou neurocientista. Sou um físico teórico com um interesse incansável pela mente. Espero que a posição vantajosa de um físico possa ajudar a enriquecer ainda mais nosso conhecimento e dar outra compreensão do objeto mais familiar e estranho do universo: nossa mente.

Mas em vista do ritmo alucinante em que estão sendo desenvolvidas novas perspectivas, é importante ter uma noção sólida da composição do cérebro. Portanto, vamos falar primeiro das origens da neurociência moderna, que alguns historiadores acreditam ter começado quando uma enorme haste de ferro atravessou o cérebro de um certo Phineas Gage. Esse acontecimento seminal gerou uma reação em cadeia que possibilitou a investigação científica do cérebro. Apesar de ter sido uma experiência infeliz para Gage, abriu o caminho para a ciência moderna.

LIVRO I A MENTE E A CONSCIÊNCIA

Minha premissa fundamental sobre o cérebro é que seu funcionamento – o que às vezes chamamos de “mente” – é consequência de sua anatomia e fisiologia, e nada mais.

– CARL SAGAN

1 DESVENDANDO A MENTE

Em 1848, Phineas Gage estava trabalhando como contramestre numa ferrovia em Vermont quando a dinamite explodiu acidentalmente, lançando uma haste de ferro de um metro e dez centímetros que atravessou a parte frontal do cérebro dele, saiu pelo alto do crânio e foi aterrissar vinte e cinco metros adiante. Os trabalhadores, horrorizados ao ver explodir uma parte do cérebro do contramestre, chamaram um médico imediatamente. Para grande espanto dos trabalhadores (e dos médicos também), Gage não morreu. Ficou semiconsciente durante semanas, mas sua recuperação, aparentemente, foi total. (Uma rara foto de Gage surgiu em 2009, mostrando um homem bonito de ar confiante, com um ferimento na cabeça e no olho esquerdo, segurando a haste de ferro.) Seus companheiros de trabalho, porém, notaram uma mudança brutal em sua personalidade depois do acidente. Normalmente alegre, prestativo, Gage ficou agressivo, hostil e egoísta. As moças eram avisadas para manter distância do contramestre. O dr. John Harlow, médico que tratou dele, observou que Gage estava “cheio de caprichos, hesitante, planejando muitas operações que eram logo abandonadas em troca de outras que pareciam mais viáveis. Uma criança em termos de capacidade intelectual e ações, mas com as paixões animais de um homem forte”. O dr. Harlow notou que ele estava “radicalmente mudado” e os companheiros de trabalho diziam que “ele não era mais o mesmo Gage”. Depois da morte de Gage em 1860, o dr. Harlow preservou seu crânio e a haste que o tinha danificado. Radiografias detalhadas do crânio confirmaram que o ferro havia causado uma grande destruição na área do cérebro atrás da testa, conhecida como lobo frontal, tanto no hemisfério esquerdo quanto no direito.

Esse incrível acidente veio mudar não só a vida de Phineas Gage, mas também alterar o curso da ciência. Anteriormente, o pensamento dominante dizia que o cérebro e a alma eram duas entidades separadas, uma filosofia chamada dualismo, mas ficou cada vez mais claro que o dano no lobo frontal havia causado mudanças abruptas na personalidade de Gage. Isso, por sua vez, criou uma mudança de paradigma no pensamento científico: talvez áreas específicas do cérebro pudessem estar relacionadas a certos comportamentos.

O CÉREBRO DE BROCA

Em 1861, apenas um ano após a morte de Gage, essa ideia foi consolidada pelo trabalho de Pierre Paul Broca, um físico de Paris que documentou um paciente que parecia normal, exceto por um grave déficit de fala. O paciente compreendia perfeitamente a fala, mas só conseguia emitir um som, “tam”.

Quando esse paciente morreu, a autópsia mostrou que ele tinha uma lesão no lobo temporal esquerdo, uma região perto da orelha esquerda. Mais tarde, o dr. Broca confirmou doze casos semelhantes de pacientes com lesão nessa área específica do cérebro. Hoje, diz-se que os pacientes com lesão no lobo temporal, geralmente do lado esquerdo, sofrem de afasia de Broca. (Em geral, os pacientes com esse distúrbio podem entender a fala mas não conseguem falar nada, ou pulam muitas palavras quando falam.)

Pouco depois, em 1874, o físico alemão Carl Wernicke descreveu pacientes que tinham o problema oposto. Eles sabiam pronunciar perfeitamente as palavras, mas não entendiam o significado, fossem faladas ou escritas. Muitas vezes, falavam fluentemente, com gramática e sintaxe corretas, mas as palavras não faziam sentido e as frases eram disparatadas. Infelizmente, quase nunca sabiam que estavam produzindo um discurso absurdo. Wernicke confirmou, depois da autópsia, que esses pacientes tinham sofrido lesões numa área ligeiramente diferente no lobo temporal esquerdo.

Os trabalhos de Broca e Wernicke representaram um marco para os estudos da neurociência, estabelecendo uma conexão clara entre problemas comportamentais, como distúrbios da fala e linguagem, e lesões em regiões específicas do cérebro.

Outra descoberta ocorreu em meio ao caos da guerra. Ao longo da história, houve muitos tabus religiosos proibindo a dissecação do corpo humano, o que restringia severamente o progresso da medicina. Na guerra, porém, com dezenas de milhares de soldados sangrando e morrendo nos campos de batalha, tornou-se uma missão urgente dos médicos desenvolver qualquer tratamento que funcionasse. Em 1864, durante a guerra da Prússia com a Dinamarca, o médico alemão Gustav Fritsch tratou de muitos soldados com ferimentos que deixavam o cérebro à mostra e observou que, quando ele tocava num hemisfério cerebral, o lado oposto do corpo estremecia. Mais tarde Fritsch demonstrou sistematicamente que, quando ele estimulava eletricamente o cérebro, o hemisfério esquerdo controlava o lado direito do corpo e vice-versa. Foi uma descoberta fantástica, demonstrando que o cérebro é de natureza basicamente elétrica, e que uma região controla o lado oposto do corpo. (Curiosamente, o uso de sondas elétricas no cérebro foi registrado alguns milhares de anos antes, pelos romanos. No ano 43 d.C. há registros de que o médico da corte do imperador Cláudio usava peixes elétricos aplicados diretamente na cabeça de pacientes que sofriam de cefaleia grave.)

A descoberta de que havia caminhos elétricos ligando o cérebro ao corpo só foi analisada sistematicamente nos anos 1930, quando o dr. Wilder Penfield começou a trabalhar com pacientes epiléticos, que geralmente sofriam convulsões e ataques, correndo risco de morrer. Para eles, a última opção era uma cirurgia cerebral, que envolvia abrir partes do crânio, expondo o cérebro.

(Como o cérebro não tem nervos sensíveis à dor, a pessoa podia ficar consciente durante todo o procedimento, e o dr. Penfield aplicava apenas uma anestesia local durante a operação.)

O médico notou que, quando estimulava certas áreas do córtex com um eletrodo, diferentes partes do corpo reagiam. Ele percebeu que poderia traçar uma correspondência direta entre regiões específicas do córtex e partes do corpo humano, e o fez com tanta precisão que é usada até hoje, quase sem alterações. Isso causou um impacto imediato, tanto na comunidade científica como no público em geral. Em um dos diagramas a seguir, pode-se ver qual região cerebral controla basicamente qual função, e quão importante é cada função. Por exemplo: como as mãos e a boca são vitais para nossa sobrevivência, uma quantidade considerável de energia cerebral é destinada para controlá-las, ao passo que nossos sensores nas costas mal são registrados.

Além disso, Penfield descobriu que ao estimular partes do lobo temporal os pacientes reativavam lembranças há muito tempo esquecidas, com total clareza. Ele ficou chocado quando, no meio de uma cirurgia cerebral, um paciente falou de repente: “Parecia... que eu estava na porta da [minha] escola secundária... Ouvi minha mãe falando ao telefone, dizendo à minha tia que fosse lá em casa à noite.” Penfield viu que estava tocando em memórias enterradas no fundo do cérebro. Quando ele publicou seus resultados, em 1951, foi criada mais uma transformação em nosso entendimento do cérebro.

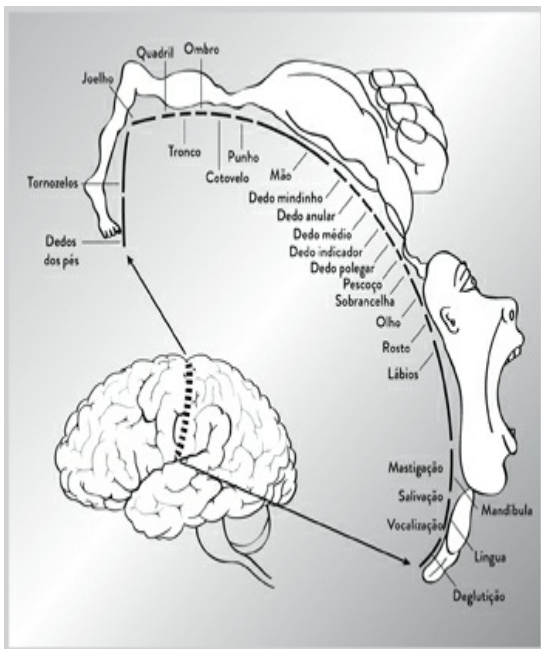


Figura 1. Mapa do córtex motor criado pelo dr. Wilder Penfield, mostrando qual região do cérebro controla qual parte do corpo.

UM MAPA DO CÉREBRO

Nos anos 1950 e 60 foi possível criar um mapa cerebral muito simples, localizando as regiões e identificando as funções de algumas delas.

Na Figura 2, vemos o neocórtex, que é a camada exterior do cérebro, dividido em quatro lobos. Essa parte é altamente desenvolvida nos humanos. Todos os lobos cerebrais são destinados a processar os sinais de nossos sentidos, exceto um: o lobo frontal, localizado atrás da testa. O córtex pré-frontal, a parte dianteira do lobo frontal, é onde quase todo o pensamento racional é processado. A informação que você está lendo agora está sendo processada em seu córtex pré-frontal. Uma lesão nessa área pode prejudicar sua capacidade de planejar ou de imaginar o futuro, como no caso de Phineas Gage. Nessa região é avaliada a informação que nossos sentidos enviam, e é decidida a ação seguinte.

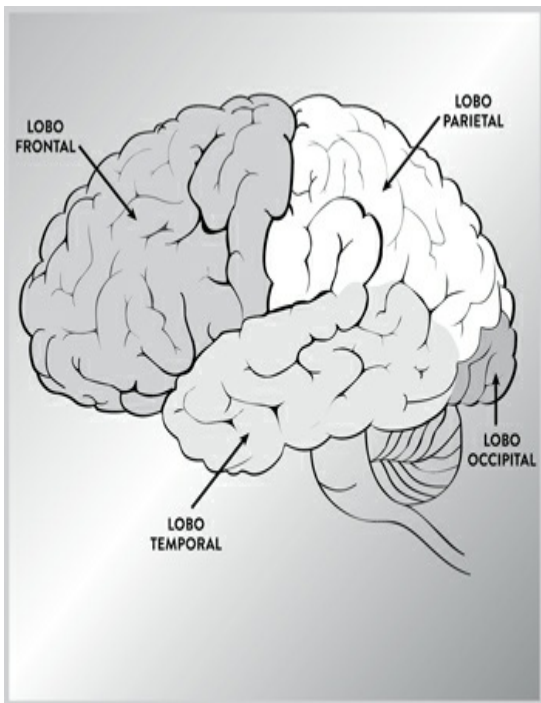


Figura 2. Os quatro lobos do neocórtex do cérebro são responsáveis por funções diferentes, embora relacionadas.

O lobo parietal está localizado no alto do cérebro. O hemisfério direito controla a atenção sensorial e a imagem corporal. O hemisfério esquerdo controla os movimentos hábeis e alguns aspectos da linguagem. Uma lesão nessa área pode

causar muitos problemas, tais como a dificuldade de localizar partes do próprio corpo.

O lobo occipital se situa na parte mais posterior do cérebro e processa a informação visual, enviada pelos olhos. Uma lesão nessa área pode causar cegueira e deficiências visuais.

O lobo temporal controla a linguagem (apenas o lado esquerdo), bem como o reconhecimento visual de rostos e certos sentimentos. Uma lesão nesse lobo pode deixar a pessoa incapaz de falar ou de reconhecer rostos familiares.

O CÉREBRO EM DESENVOLVIMENTO

Quando olhamos os outros órgãos do corpo, como músculos, ossos e pulmões, vemos imediatamente que eles têm uma constituição óbvia, mas a estrutura do cérebro pode parecer um emaranhado caótico. De fato, as tentativas de mapeá-lo foram muitas vezes chamadas de “cartografia para tolos”.

Para dar sentido à estrutura aparentemente aleatória do cérebro, em 1967 o dr. Paul MacLean, do National Institute of Mental Health, aplicou ao cérebro a teoria da evolução proposta por Charles Darwin. Ele dividiu o cérebro em três partes. (Desde então, outros estudos acrescentaram muitos detalhes, mas usaremos esse modelo como um princípio básico de organização para explicar a estrutura geral do cérebro.) Primeiro, ele observou que as partes central e posterior, contendo o tronco encefálico, o cerebelo e os gânglios basais, são quase idênticas às do cérebro dos répteis. Conhecidas como constituintes do “cérebro reptiliano”, elas são as estruturas cerebrais mais antigas, governando as funções animais básicas, como equilíbrio, respiração, digestão, batimentos cardíacos e pressão sanguínea. Elas controlam também comportamentos como enfrentamento, caça, acasalamento e territorialidade, necessários para a sobrevivência e reprodução. O cérebro reptiliano remonta a cerca de 500 milhões de anos. (Ver Figura 3.)

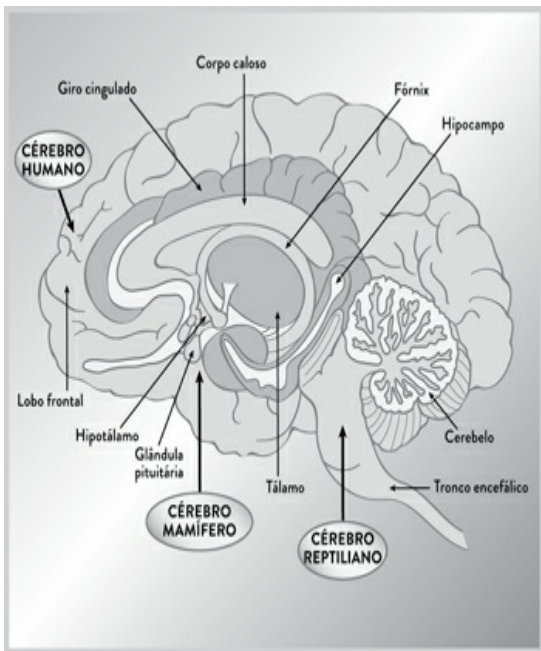


Figura 3. A história evolucionária cerebral, com o cérebro reptiliano, o sistema límbico (o cérebro dos mamíferos) e o neocórtex (o cérebro humano). De um modo geral, pode-se dizer que o caminho da evolução do nosso cérebro passou do reptiliano para o mamífero e daí para o humano.

À medida que evoluímos de répteis para mamíferos, o cérebro foi se tornando mais complexo, expandindo-se e criando estruturas inteiramente novas. Formou-se assim o “cérebro mamífero”, ou sistema límbico, localizado perto do centro do cérebro, rodeando partes do reptiliano. O sistema límbico é proeminente em

animais que vivem em grupos sociais, como os macacos. Também possui estruturas relacionadas com as emoções. Como a dinâmica dos grupos sociais pode ser muito complexa, o sistema límbico é essencial para identificar possíveis inimigos, aliados e rivais.

As partes do sistema límbico que controlam os comportamentos cruciais para o convívio de animais sociais são:

- O hipocampo. É o portal da memória, onde lembranças recentes são processadas para virar memórias de longo prazo. O nome significa “cavalo-marinho”, devido a sua forma peculiar. Uma lesão nessa região destrói a capacidade de guardar memórias antigas. Você se torna prisioneiro do presente.
- A amígdala. É o local das emoções, principalmente do medo. É lá que as emoções são primeiro registradas e geradas. Seu nome significa “amêndoa”.
- O tálamo. É como uma estação de transmissão, que recolhe sinais do tronco encefálico e os envia aos vários níveis corticais. Seu nome quer dizer “câmara interna”.
- O hipotálamo. Regula a temperatura do corpo, o ritmo circadiano, a fome, a sede, e aspectos da reprodução e do prazer. Situa-se abaixo do tálamo – daí seu nome.

Por fim, temos a terceira e mais recente região do cérebro mamífero, o córtex cerebral, que é a camada externa do cérebro. A última estrutura cerebral a surgir da evolução é o neocórtex (que significa “casca nova”). Tal estrutura governa o comportamento cognitivo superior, sendo mais desenvolvida em humanos. Ocupa 80% da nossa massa cerebral, apesar de ser fina como um guardanapo. Nos ratos, o neocórtex é liso, mas nos humanos é cheio de circunvoluções, permitindo que uma grande superfície se acomode dentro do crânio.

Em certo sentido, o cérebro humano é como um museu, guardando resquícios dos estágios anteriores de nossa evolução ao longo de milhões de anos, expandindo para fora e para frente em tamanho e função. (Este é também o caminho, grosso modo, trilhado por uma criança recém-nascida. O cérebro do bebê se expande para fora e para frente, talvez reproduzindo os estágios de nossa evolução.)

Apesar de o neocórtex parecer simples, as aparências enganam. Ao microscópio, podemos apreciar a complexa arquitetura do cérebro. A massa cinzenta consiste em bilhões de células minúsculas chamadas neurônios. Como

uma gigantesca rede telefônica, eles recebem mensagens de outros neurônios por meio dos dendritos, que são como raízes brotando em uma extremidade do neurônio. Na outra extremidade há uma fibra longa, chamada axônio. O axônio conecta-se a até dez mil outros neurônios através dos dendritos. Na junção entre os dois há uma brecha minúscula, chamada sinapse. Essas sinapses agem como portões, regulando o fluxo de informações dentro do cérebro. Substâncias químicas especiais, chamadas neurotransmissores, podem entrar na sinapse e alterar o fluxo de sinais. Os neurotransmissores – como a dopamina, a serotonina e a noradrenalina – ajudam a controlar a corrente de informações que se movimentam por miríades de percursos neurais, e por isso exercem um efeito poderoso sobre nosso humor, nossas emoções, e nossos pensamentos e estado mental. (Ver Figura 4.)

Essa descrição do cérebro era a representação básica do nível de conhecimento na década de 1980. Nos anos 1990, porém, com a introdução de novas tecnologias no campo da física, o mecanismo do pensamento passou a ser revelado com muito mais detalhes, desencadeando a explosão atual de descobertas científicas. Um dos motores dessa revolução foi o aparelho de IRM.

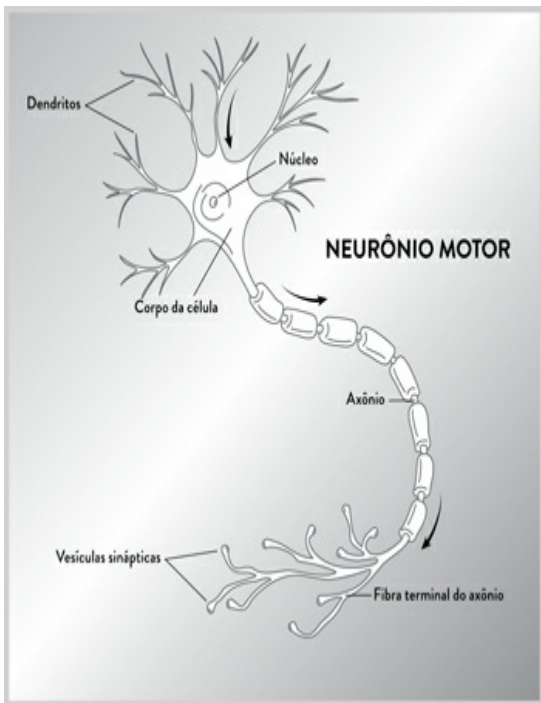


Figura 4. Diagrama de um neurônio. Sinais elétricos viajam ao longo do axônio até encontrar a sinapse. Os neurotransmissores podem regular o fluxo de sinais que passam pela sinapse.

IMAGEM POR RESSONÂNCIA MAGNÉTICA (IRM): UMA JANELA PARA O CÉREBRO

Para entender as razões pelas quais essa nova e fundamental tecnologia ajudou a decodificar o cérebro pensante, é preciso retomar alguns princípios básicos da física.

As ondas de rádio, que são um tipo de radiação eletromagnética, podem atravessar diretamente o tecido sem causar dano. Os aparelhos de IRM tiram proveito disso, permitindo que as ondas eletromagnéticas penetrem livremente no crânio. Assim, essa tecnologia tem apresentado fotografias sensacionais de algo antes considerado impossível de ser captado: o funcionamento interno do cérebro enquanto experimenta sensações e emoções. Ao observar a dança das luzes piscando no aparelho de IRM, podemos acompanhar os pensamentos se movimentando dentro do cérebro. É como poder ver dentro de um relógio em funcionamento.

O primeiro aspecto que se nota num aparelho de IRM é o enorme cilindro da bobina magnética, capaz de produzir um campo magnético de 20 a 60 mil vezes maior que o da Terra. Esse ímã gigantesco é um dos principais motivos para que um aparelho de IRM pese uma tonelada, ocupe uma sala inteira e custe muitos milhões de dólares. (Os aparelhos de IRM são mais seguros que os de raios X porque não criam íons perigosos. A tomografia computadorizada [TC], que também pode criar imagens em 3D, expõe o corpo a uma dose muito maior do que o raio X comum, e por isso tem que ser muito bem regulada. Em contraste, as máquinas de IRM são seguras quando usadas adequadamente. Um problema, porém, é o descuido de operadores. O campo magnético é forte o suficiente para arremessar instrumentos pelos ares, se o aparelho for ligado no momento errado. Muitos já se machucaram, e houve até mortes.)

Os aparelhos de IRM funcionam assim: o paciente se deita e é inserido num cilindro contendo duas grandes bobinas, que criam o campo magnético. Quando o campo magnético é ligado, os núcleos dos átomos dentro do corpo agem como uma agulha de bússola, alinhando-se horizontalmente na direção do campo. Então é gerado um pequeno pulso de frequência de rádio, fazendo com que alguns núcleos do corpo se virem na direção oposta. Quando os núcleos voltam à posição anterior, emitem um pulso secundário de rádio, que é então analisado pelo aparelho de IRM. A análise desses pequenos “ecos” permite reconstruir a localização e a natureza desses átomos. Como o morcego, que usa o som para determinar a posição dos objetos em seu caminho, os ecos criados pela IRM permitem aos cientistas recriar uma imagem extraordinária do interior do cérebro. Os computadores reconstroem a posição dos átomos, fornecendo belos diagramas em três dimensões.

Quando foi introduzida, a IRM mostrava a estrutura estática das várias regiões

cerebrais. Entretanto, em meados dos anos 1990, foi inventado um novo tipo de IRM, chamado “funcional”, ou IRMf, que detectava a presença de oxigênio no sangue do cérebro. (Para diferentes tipos de IRM, os cientistas colocam uma letra minúscula na frente, mas vamos manter a mesma sigla para designar todos os tipos de aparelho IRM.) O IRM não consegue detectar diretamente o fluxo de eletricidade nos neurônios, mas como o oxigênio é necessário para fornecer energia aos neurônios, o sangue oxigenado pode rastrear indiretamente o fluxo da energia elétrica nos neurônios e mostrar como as várias regiões do cérebro interagem umas com as outras.

A varredura por IRM já invalidou definitivamente a ideia de que o pensamento está concentrado num único ponto. Em vez disso, pode-se ver a energia elétrica circulando em partes diferentes enquanto o cérebro pensa. Acompanhando o percurso dos pensamentos, a varredura por IRM lançou uma nova luz sobre a natureza dos males de Alzheimer e Parkinson, da esquizofrenia e de diversas outras doenças mentais.

A grande vantagem das máquinas de IRM é a sofisticada capacidade de localizar partes do cérebro tão minúsculas quanto uma fração de milímetro. Uma varredura por IRM não cria somente pontos numa tela bidimensional, chamados pixels, mas também pontos no espaço tridimensional, chamados “vóxeis”, produzindo um conjunto luminoso de dezenas de milhares de pontos coloridos em 3D, no formato de um cérebro.

Como cada elemento químico responde a frequências de rádio diferentes, pode-se alterar a frequência do pulso de rádio e assim identificar diversos elementos do corpo. Como foi dito, aparelhos de IRMf concentram-se nos átomos de oxigênio contidos no sangue a fim de medir o fluxo sanguíneo, mas os aparelhos de IRM podem também ser sintonizados para identificar outros átomos. Apenas na última década foi introduzida uma nova forma de IRM chamada de “imagem de tensor de difusão” (ITD), que detecta o fluxo de água no cérebro. Como a água segue os caminhos neurais, o aparelho de ITD produz belas figuras, que parecem trepadeiras entrelaçadas crescendo num jardim. Assim os cientistas podem determinar instantaneamente como certas partes do cérebro estão ligadas a outras.

Entretanto, a tecnologia de IRM tem algumas desvantagens. Embora tenha uma qualidade incomparável em termos de resolução espacial, sendo capaz de mostrar vóxeis do tamanho de uma cabeça de alfinete em três dimensões, não é tão boa em resolução temporal. Leva quase um segundo inteiro para seguir o percurso do sangue no cérebro, o que pode parecer muito pouco, mas os sinais elétricos viajam pelo cérebro quase instantaneamente e, portanto, a IRM pode perder alguns dos detalhes mais complexos dos padrões de pensamento.

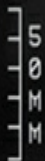
Outra dificuldade é o preço, que chega a milhões de dólares, de modo que os médicos precisam compartilhar o aparelho. Mas, assim como muitas tecnologias

e seus aprimoramentos, ao longo do tempo o preço deve baixar.

O custo exorbitante não deteve a busca por aplicações comerciais. Uma ideia é usar a IRM como detector de mentiras, pois, segundo alguns estudos, pode identificar mentiras com 95%, ou mais, de certeza. O nível de certeza ainda é controverso, mas a ideia básica é que, quando alguém mente, tem que simultaneamente saber a verdade, inventar a mentira e analisar rapidamente a coerência da mentira com os fatos conhecidos. Atualmente, algumas empresas afirmam que a tecnologia de IRM mostra que os lobos pré-frontal e parietal se iluminam quando alguém conta uma mentira. Mais especificamente, é ativado o “córtex orbitofrontal” (que, entre outras funções, serve como “verificador de fatos” para nos avisar quando algo está errado). Essa área se localiza logo atrás das órbitas dos olhos, daí seu nome. Diz a teoria que o córtex orbitofrontal entende a diferença entre verdade e mentira, ficando hiperativo quando detecta uma mentira. (Outras áreas do cérebro também acendem quando alguém conta uma mentira, como os córtices pré-frontal superior medial e inferolateral, relacionados com a cognição.)

Várias empresas já estão oferecendo aparelhos de IRM como detectores de mentiras, que começam a ser utilizados em casos judiciais. Mas é importante observar que as varreduras por IRM só indicam aumento de atividade cerebral em certas áreas. Enquanto exames de DNA podem mostrar com precisão uma parte em 10 bilhões ou mais, na varredura por IRM isso não acontece, porque é preciso ativar várias áreas do cérebro para se inventar uma mentira, e essas mesmas áreas são responsáveis também pelo processamento de outros tipos de pensamento.

SETON
15-FEB
19
SCAN



P U 1
L

GYROSCAN T5-II

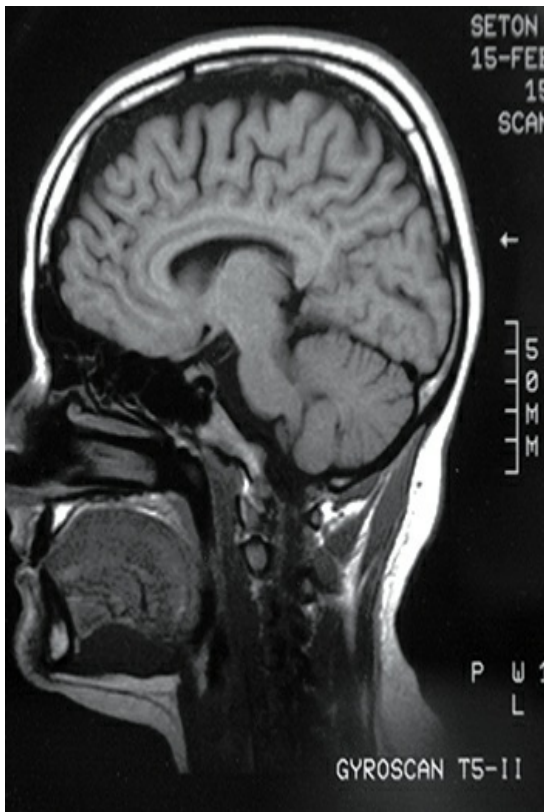




Figura 5. Na primeira figura, vemos uma imagem fornecida por um aparelho de IRM funcional mostrando regiões de alta atividade mental. Na segunda, vemos a imagem típica em forma de flor, criada por um aparelho de IRM de difusão, capaz de acompanhar os percursos neurais e as conexões cerebrais.

VARREDURA POR ELETROENCEFALOGRAMA (EEG)

Outra ferramenta que investiga o cérebro a fundo é o EEG, o eletroencefalograma. O EEG foi introduzido há muito tempo, em 1924, mas só recentemente foi possível empregar computadores para dar sentido a todos os dados fornecidos por eletrodo.

Para usar o aparelho de EEG, em geral o paciente coloca um capacete de aparência futurista com vários eletrodos na superfície. (Versões mais avançadas colocam na cabeça do paciente uma redinha contendo uma série de pequenos eletrodos.) Esses eletrodos detectam os pequenos sinais elétricos que circulam no cérebro.

Uma varredura por EEG difere de uma por IRM em vários aspectos cruciais. Como vimos, a varredura por IRM emite pulsos de rádio no cérebro e depois analisa os “ecos” que retornam. Isso significa que é possível variar o pulso de rádio a fim de selecionar diferentes átomos para análise, tornando o aparelho muito versátil. O aparelho de EEG, porém, é estritamente passivo, isto é, ele analisa os pequenos sinais eletromagnéticos que o cérebro emite naturalmente. O EEG é excelente para registrar os vastos sinais eletromagnéticos que irrompem em todo o cérebro, o que permite aos cientistas medir a capacidade geral do cérebro quando dorme, quando se concentra, relaxa, sonha etc. Diferentes estados de consciência vibram em diferentes frequências. Por exemplo: o sono profundo corresponde a ondas delta, que vibram de 0,1 a 4 ciclos por segundo. Estados mentais ativos, como na resolução de problemas, correspondem a ondas beta, vibrando de 12 a 30 ciclos por segundo. Essas vibrações permitem que várias áreas do cérebro compartilhem informações e se comuniquem umas com as outras, mesmo estando localizadas em lados opostos. Além disso, enquanto o fluxo de sangue por IRM só pode ser medido algumas vezes por segundo, o EEG mede a atividade elétrica instantaneamente.

As maiores vantagens do EEG, porém, são a conveniência e o custo. Até alunos do ensino médio fizeram experimentos em suas casas com sensores de EEG colocados na cabeça.

Entretanto, a grande desvantagem do EEG, que impediu seu desenvolvimento durante décadas, é a grande deficiência em termos de resolução espacial. O EEG capta sinais elétricos que já foram difundidos depois de passar pelo crânio, tornando difícil detectar atividades anormais originadas nas profundezas do cérebro. Observando a produção dos confusos sinais do EEG, é quase impossível dizer com certeza em qual parte do cérebro eles foram criados. Além disso, movimentos lentos, como mexer um dedo, podem distorcer o sinal, às vezes tornando o procedimento inútil.

VARREDURAS POR TOMOGRAFIA POR EMISSÃO DE PÓSITRONS (TEP)

Outro instrumento do mundo da física é a tomografia por emissão de pósitrons (TEP), que calcula o fluxo de energia no cérebro localizando a presença de glicose, a molécula de açúcar que alimenta as células. Assim como a câmara de

nuvens que fiz quando estudante, a varredura por TEP utiliza partículas subatômicas emitidas pelo sódio-22 dentro da glicose. Para dar início à varredura por TEP, injeta-se no paciente uma solução especial contendo um açúcar levemente radioativo. Os átomos de sódio dentro das moléculas de açúcar são substituídos por átomos de sódio-22 radioativo. Cada vez que um átomo de sódio decai, emite um elétron positivo, ou pósitron, que é facilmente detectado pelos sensores. Acompanhando o percurso dos átomos de sódio radioativo no açúcar, é possível traçar o fluxo de energia dentro do cérebro vivo.

Em vários aspectos, a varredura por TEP consegue os mesmos benefícios que a por ressonância magnética, mas não tem a resolução espacial detalhada de uma foto de IRM. Contudo, em vez de medir o fluxo do sangue, que é apenas um indicador indireto do consumo de energia no corpo, a varredura por TEP mede o consumo de energia, portanto está mais ligada à atividade neural.

Mas há outra desvantagem. Ao contrário da IRM e do EEG, a varredura por TEP é levemente radioativa, e desse modo os pacientes não podem ser submetidos a ela de forma contínua. Em geral, a varredura por TEP não pode ser aplicada a uma pessoa mais que uma vez por ano, devido aos riscos da radiação.

MAGNETISMO NO CÉREBRO

Na última década, vários novos aparelhos de alta tecnologia passaram a ser utilizados por cientistas, inclusive a varredura eletromagnética transcraniana (EMT), a magnetoencefalografia (MEG), a espectroscopia no infravermelho próximo (NIRS, da sigla em inglês para *near-infrared spectroscopy*) e a optogenética, entre outros.

O magnetismo tem sido usado para desligar, de forma sistemática, partes específicas do cérebro sem a necessidade de abri-lo. A física básica por trás desses novos instrumentos é que um campo elétrico de rápida variação pode criar um campo magnético e vice-versa. A MEG mede passivamente os campos magnéticos produzidos pelas variações dos campos elétricos do cérebro. Esses campos magnéticos são fracos e extremamente pequenos, apenas um bilionésimo do campo magnético da Terra. Assim como o EEG, a MEG é muito boa na resolução temporal, chegando a um milésimo de segundo. Sua resolução espacial, porém, é de apenas um centímetro cúbico.

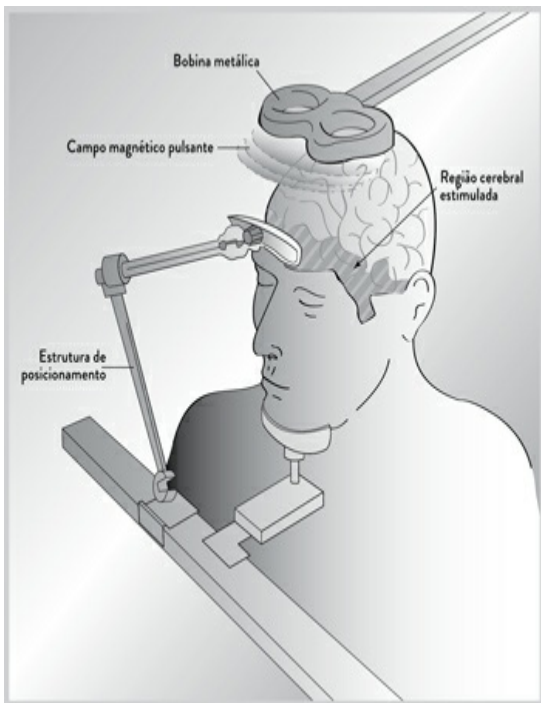


Figura 6. Vemos o aparelho de EMT e o magnetoencefalógrafo, que utiliza o magnetismo, em vez de ondas de rádio, para penetrar no crânio e determinar a natureza dos pensamentos no interior do cérebro. O magnetismo pode silenciar temporariamente partes do cérebro, permitindo aos cientistas determinar com segurança o desempenho dessas regiões, sem depender de vítimas de derrame.

Ao contrário da medição passiva da MEG, a EMT gera um grande pulso de eletricidade, que cria um surto de energia magnética. O aparelho de EMT é colocado próximo ao cérebro, de modo que o pulso magnético penetra no crânio e cria outro pulso elétrico no interior do cérebro. Esse pulso elétrico secundário, por sua vez, é suficiente para desligar ou amortecer a atividade das áreas cerebrais selecionadas.

Historicamente, os cientistas dependiam da ocorrência de derrames ou tumores para silenciar certas partes do cérebro, e então determinar o que elas faziam. Com a EMT podem-se desligar ou amortecer partes do cérebro, em segurança, conforme a necessidade. Aplicando energia magnética num determinado local do cérebro, é possível definir sua respectiva função por meio da simples observação de mudanças de comportamento do paciente. (Por exemplo: aplicando pulsos magnéticos no lobo temporal esquerdo, observa-se que a capacidade de falar é comprometida.)

Uma possível desvantagem da EMT é que esses campos magnéticos não penetram muito fundo no interior do cérebro (porque os campos magnéticos diminuem muito mais rapidamente do que a usual lei do inverso do quadrado da distância, para a eletricidade). A EMT é muito útil para desligar partes do cérebro perto do crânio, mas o campo magnético não atinge centros importantes em locais mais profundos, como o sistema límbico. Contudo, as futuras gerações de aparelhos de EMT podem vir a superar esse problema técnico aumentando a intensidade e a precisão do campo magnético.

ESTIMULAÇÃO CEREBRAL PROFUNDA

Outra ferramenta que tem sido vital para os neurologistas é a estimulação cerebral profunda (ECP). As sondagens inicialmente realizadas pelo dr. Penfield eram relativamente toscas. Hoje os eletrodos podem ser finos como fios de cabelos e atingir áreas específicas no interior mais profundo do cérebro. A ECP não só permite localizar a função de várias partes do cérebro, como também pode ser usada para tratar distúrbios mentais. A ECP já provou o seu valor no tratamento da doença de Parkinson, em que certas regiões cerebrais ficam hiperativas e costumam gerar um tremor incontrollável das mãos.

Mais recentemente, esses eletrodos tiveram como alvo uma nova área cerebral (chamada área de Brodmann número 25), que costuma ficar ativa demais em pacientes deprimidos que não reagem nem a psicoterapias nem a medicamentos. A estimulação cerebral profunda tem dado um alívio miraculoso, após décadas de tormento e agonia, a pacientes que sofrem de depressão há muitos anos.

A cada ano são encontradas novas finalidades para a estimulação cerebral profunda. De fato, praticamente todas as principais doenças do cérebro têm sido reexaminadas à luz dessa e de outras novas tecnologias de varredura cerebral. Isso promete ser um novo campo importante para o diagnóstico, e até para o tratamento das doenças.

OPTOGENÉTICA – ILUMINANDO O CÉREBRO

Talvez o mais novo e animador recurso dos neurologistas seja a optogenética, antes considerada ficção científica. Como uma varinha mágica, emitindo um fecho de luz no cérebro, permite ativar certos caminhos que controlam o comportamento.

Inacreditavelmente, é possível inserir, com precisão cirúrgica, um gene sensível à luz diretamente num neurônio. Então, acendendo-se uma luz, o neurônio é ativado. O mais importante é que os cientistas conseguem acessar esses percursos, podendo ligar e desligar certos comportamentos por meio de um interruptor.

Apesar de ter apenas uma década, a tecnologia da optogenética já se mostrou eficaz no controle de certos comportamentos de animais. Acionando um interruptor de luz, é possível fazer moscas de fruta saírem voando, minhocas pararem de se contorcer e camundongos correrem loucamente em círculos. Foram iniciados experimentos com macacos, e já estão em discussão experimentos com humanos. Há grande esperança de que essa tecnologia tenha uma aplicação direta no tratamento de distúrbios como o Parkinson e depressão.

O CÉREBRO TRANSPARENTE

Outra descoberta espetacular torna o cérebro totalmente transparente, de modo que os caminhos neurais podem ser vistos a olho nu. Em 2013, cientistas da Universidade de Stanford anunciaram que conseguiram tornar transparente o cérebro inteiro de um camundongo, e também partes de um cérebro humano. O anúncio foi tão espantoso que ganhou a primeira página do *New York Times*, com a manchete “Cérebro transparente como gelatina para a investigação científica”.

No nível celular, as células vistas individualmente são transparentes, com exposição total de todos os seus componentes microscópicos. Porém, quando bilhões de células se unem para formar órgãos como o cérebro, os lipídios agregados (gorduras, óleos, ceras e substâncias químicas não solúveis em água)

tornam o órgão opaco. A chave da nova técnica é remover os lipídios, mas deixar os neurônios intactos. Os cientistas de Stanford fizeram isso colocando o cérebro em hidrogel (uma substância como o gel, composta principalmente de água), que se liga a todas as moléculas, exceto aos lipídios. Colocando o cérebro numa solução saponácea, com um campo elétrico, os lipídios se integram à substância, que quando é descartada deixa o cérebro transparente. A adição de corantes torna visíveis os caminhos neurais. Isso possibilita identificar e mapear os diversos caminhos neurais do cérebro.

Tornar um tecido transparente não é novidade, mas conseguir as condições necessárias para deixar o cérebro inteiro transparente exigiu muita criatividade. “Queimei e derreti mais de cem cérebros”, confessa o dr. Kwanghun Chung, um dos principais cientistas por trás do estudo. A nova técnica, chamada *Clarity*, pode ser aplicada também a outros órgãos (até órgãos preservados durante anos em substâncias químicas como o formol). Ele já criou fígados, pulmões e corações transparentes. Essa nova técnica tem aplicações admiráveis em toda a medicina. Mas, principalmente, irá acelerar a localização de caminhos neurais no cérebro, que é o foco de intensas pesquisas e financiamentos.

QUATRO FORÇAS FUNDAMENTAIS

O sucesso dessa primeira geração de varreduras cerebrais é espetacular. Antes de sua introdução, apenas cerca de trinta regiões cerebrais eram conhecidas com alguma certeza. Hoje, o aparelho de IRM sozinho consegue identificar de duzentas a trezentas regiões, abrindo fronteiras inteiramente novas para a ciência do cérebro. Em vista de tantas novas tecnologias de varredura introduzidas pela física apenas nos últimos quinze anos, cabe perguntar: Haverá mais? A resposta é sim, mas serão variações e refinamentos das anteriores, e não tecnologias radicalmente novas. Isso porque só existem quatro forças fundamentais – gravitacional, eletromagnética, nuclear fraca e nuclear forte – que regem o universo. (Os físicos vêm tentando encontrar evidências de uma quinta força, mas até o momento essas tentativas fracassaram.)

A força eletromagnética, que ilumina as cidades e representa a energia da eletricidade e do magnetismo, é a fonte de quase todas as novas tecnologias de varredura (à exceção da TEP, que é regida pela força nuclear fraca). Uma vez que os físicos têm mais de 150 anos de experiência no uso da força eletromagnética, não há mistério em criar novos campos elétricos e magnéticos, portanto, qualquer nova tecnologia de varredura cerebral será muito provavelmente uma modificação das tecnologias existentes, e não algo inteiramente novo. Assim como ocorre com muitas tecnologias, o tamanho e o

custo dos aparelhos irão diminuir, aumentando significativamente a utilização dessas sofisticadas máquinas. Os físicos já estão fazendo os cálculos necessários para fazer um aparelho de IRM caber num telefone celular. Ao mesmo tempo, o principal desafio para as varreduras cerebrais é a resolução, tanto espacial como temporal. A resolução espacial da IRM irá aumentar na medida em que o campo magnético se torne mais uniforme e que os componentes eletrônicos fiquem mais sensíveis. Atualmente, a IRM só vê pontos ou vóxeis dentro duma fração de milímetro. Mas cada ponto pode conter centenas de milhares de neurônios. As novas tecnologias de varredura devem reduzir ainda mais esse espaço. O santo graal dessa abordagem será criar um aparelho do tipo IRM que possa identificar cada neurônio e suas conexões.

A resolução temporal de aparelhos de IRM é limitada também porque eles analisam o fluxo do sangue oxigenado no cérebro. O aparelho propriamente dito tem uma resolução temporal muito boa, mas rastrear o fluxo do sangue diminui sua velocidade. No futuro, outros aparelhos de IRM serão capazes de localizar diferentes substâncias mais diretamente ligadas à ativação dos neurônios, permitindo assim uma análise em tempo real dos processos mentais. Por mais espetaculares que tenham sido os progressos nos últimos quinze anos, eles são apenas uma amostra do futuro.

NOVOS MODELOS DO CÉREBRO

Historicamente, a cada descoberta científica surge um novo modelo do cérebro. Um dos mais antigos é o do “homúnculo”, um homenzinho que vivia dentro do cérebro e tomava todas as decisões. Essa imagem não era muito útil, porque não dizia o que se passava no cérebro do homúnculo. Talvez houvesse um homúnculo oculto dentro do homúnculo.

A introdução de aparelhos mecânicos simples trouxe outro modelo do cérebro: o de uma máquina, como um relógio, com rodinhas e engrenagens. Essa analogia foi útil para cientistas e inventores como Leonardo da Vinci, que até projetou um homem mecânico.

No final dos anos 1800, quando o vapor esculpia novos impérios, surgiu outra analogia, a da máquina a vapor, com fluxos de energia competindo entre si. Os historiadores acreditavam que o modelo hidráulico tinha afetado a imagem que Sigmund Freud fazia do cérebro, onde havia uma luta contínua entre três forças: o ego (representando o eu e o pensamento racional), o id (representando os desejos reprimidos) e o superego (representando a nossa consciência). Nesse modelo, se houvesse muita pressão devido a um conflito entre os três, poderia haver uma regressão ou um colapso geral de todo o sistema. Esse modelo era

engenhoso, mas, como o próprio Freud admitiu, requeria estudos detalhados do cérebro em nível neuronal, o que levaria outro século.

No início do século XX, com a ascensão do telefone, apareceu outra analogia: a de uma mesa telefônica gigantesca. O cérebro seria uma malha de linhas telefônicas conectadas a uma vasta rede. A consciência seria uma longa fileira de telefonistas em frente a um grande painel de interruptores, continuamente ligando e desligando fios. Infelizmente, esse modelo nada dizia sobre como as mensagens eram ligadas entre si para formar o cérebro.

Com a invenção do transistor, um outro modelo entrou em voga: o computador. As velhas centrais de comutação foram substituídas por microchips contendo centenas de milhões de transistores. Talvez a “mente” fosse apenas um programa de software rodando em “*wetware*” (isto é, tecido cerebral em vez de transistores). Esse modelo permanece até hoje, mas tem limitações. O modelo transistor não explica como o cérebro efetua computações que exigiriam um computador do tamanho da cidade de Nova York. E o cérebro não tem programa, nem sistema operacional Windows, nem chip Pentium. (Além disso, um microcomputador com chip Pentium é extremamente rápido, mas tem um problema. Todos os cálculos passam por esse único processador. O cérebro é o oposto. A ativação de cada neurônio é relativamente lenta, mas isso é amplamente compensado por existirem 100 bilhões de neurônios processando dados simultaneamente. Portanto, um processador paralelo lento pode superar um único processador muito veloz.)

A analogia mais recente é com a internet, que une bilhões de computadores. Nessa imagem, a consciência é um fenômeno “emergente”, surgindo milagrosamente da ação coletiva de bilhões de neurônios. (O problema dessa imagem é que não diz absolutamente nada sobre como ocorre esse milagre. Varre toda a complexidade do cérebro para baixo do tapete da teoria do caos.)

Sem dúvida, todas essas analogias têm alguma verdade, mas nenhuma delas capta realmente a complexidade do cérebro. Entretanto, uma analogia que achei útil (embora também imperfeita) é a de uma grande corporação. Nessa analogia há uma imensa burocracia e níveis de autoridade, com vastos fluxos de informação canalizados em setores diferentes. Mas as informações importantes acabam chegando ao diretor-geral, no centro de comando. Ali são tomadas as decisões finais.

Se a analogia com uma grande corporação for válida, deve conseguir explicar certos traços peculiares do cérebro:

- **A maior parte das informações é “subconsciente”.** isto é, o diretor-geral não conhece todas as informações amplas e complexas que fluem continuamente pela burocracia. De fato, apenas uma

quantidade mínima de informações chega à mesa do diretor, que pode ser comparado ao córtex pré-frontal. O diretor só precisa das informações importantes, que chamam a sua atenção; caso contrário, ele ficaria paralisado no meio de uma avalanche de informações irrelevantes.

- Essa situação é provavelmente subproduto da evolução, pois nossos ancestrais talvez ficassem sobrecarregados com informações supérfluas, subconscientes, inundando o cérebro quando confrontados com uma emergência. Felizmente, não temos consciência dos trilhões de cálculos sendo processados em nosso cérebro. Ao encontrar um tigre na floresta, não precisamos nos preocupar com o estado do nosso estômago, dedos do pé, cabelos etc. Só precisamos saber correr.
- **“Emoções” são decisões rápidas tomadas de forma independente num nível inferior.** Dado que o pensamento racional toma muitos segundos, geralmente é impossível reagir de forma racional a uma emergência. Portanto, as regiões inferiores do cérebro precisam avaliar rapidamente a situação e optam por uma decisão, uma emoção, sem autorização de cima.
- Assim, as emoções (medo, raiva, horror etc.) são sinais instantâneos de alerta emitidos por um nível inferior, gerados pela evolução, para avisar ao centro de comando sobre situações possivelmente graves ou perigosas. Temos pouco controle consciente sobre as emoções. Por exemplo: por mais que ensaiemos uma palestra para um grande público, ainda assim ficamos nervosos.

Rita Carter, autora de *O livro de ouro da mente*, escreve que “as emoções não são sentimentos, e sim um conjunto de mecanismos de sobrevivência enraizados no corpo que evoluíram para nos livrar do perigo e nos impulsionar em direção a coisas que possam nos beneficiar”.
- **Há um constante clamor pela atenção do diretor-geral.** Não é um só homúnculo, CPU ou chip de Pentium que toma as decisões. Pelo contrário, os vários subcentros sob o comando central estão em constante competição entre si, lutando pela atenção do diretor. Assim, não há uma continuidade equilibrada de pensamento, mas uma cacofonia de vários ciclos de feedback competindo uns com os outros. O conceito do “eu” como um todo unificado, sempre tomando todas as decisões, é uma ilusão criada pela nossa própria mente subconsciente.

- Costumamos achar que nossa mente é uma entidade única, processando informações de forma contínua e tranquila, totalmente encarregada das decisões. Mas as imagens que aparecem nas varreduras cerebrais divergem dessa percepção.

Marvin Minsky, professor do MIT e um dos pais da inteligência artificial, disse-me que a mente é mais parecida com uma “sociedade de mentes”, com diversos submódulos, cada um tentando competir com os outros.

Quando entrevistei Steven Pinker, psicólogo da Universidade de Harvard, perguntei como a consciência emerge dessa confusão. Ele disse que a consciência é como uma tempestade assolando o cérebro. Ele elaborou essa imagem quando escreveu que “o sentimento intuitivo que temos de que existe um ‘eu’ diretor instalado numa sala de controle no cérebro, examinando as telas dos sentidos e apertando os botões dos músculos, é uma ilusão. A consciência consiste em um turbilhão de eventos distribuídos pelo cérebro. Esses eventos competem por atenção, e quando um processo grita mais alto que os outros, o cérebro racionaliza o resultado e gera a impressão de que um único eu estava na chefia o tempo todo”.

- **As decisões finais são tomadas pelo diretor-geral no centro de comando.** Uma das funções da burocracia é compilar e organizar as informações para o diretor-geral, que se reúne apenas com os gerentes de cada setor. O diretor tenta mediar todas as informações conflitantes despejadas no centro de comando. Aqui termina a confusão. O diretor, situado no córtex pré-frontal, tem que tomar a decisão final. Enquanto nos animais muitas decisões são tomadas por instinto, os humanos tomam decisões de uma forma mais sofisticada, após analisar diferentes conjuntos de informação dados pelos sentidos.
- **Os fluxos de informação são hierárquicos.** Devido ao volume de informações, que devem subir para a sala do diretor, ou descer para a equipe de apoio, é preciso organizá-las em complexas matrizes de redes interligadas, com várias ramificações. Pense num pinheiro, com o centro de comando no alto e a pirâmide de ramos voltados para baixo, desmembrando-se em vários subcentros.
- Certamente, há diferenças entre uma burocracia e a estrutura do pensamento. A primeira norma de qualquer burocracia é que “se expande para preencher o espaço a ela destinado”. Entretanto, gastar energia é um luxo ao qual o cérebro não pode se dar. O cérebro

consome apenas cerca de vinte watts de potência (a potência de uma lâmpada muito fraca), mas é provavelmente o máximo de energia que pode consumir sem que o corpo pare de funcionar. Se o cérebro gerar mais calor, vai danificar os tecidos. Por isso, o cérebro pega atalhos continuamente para conservar energia. Veremos neste livro os dispositivos inteligentes e engenhosos, criados pela evolução, sem nosso conhecimento, para economizar.

A “REALIDADE” É MESMO REAL?

Todo mundo conhece a expressão “ver para crer”. Mas muito do que vemos é na verdade uma ilusão. Por exemplo: quando vemos uma típica paisagem bucólica, parece uma linda cena de cinema. Na realidade, existe uma lacuna, um furo, no nosso campo de visão, que corresponde à localização do nervo óptico na retina. Deveríamos enxergar esse ponto negro grande e feio em tudo o que vemos. No entanto, nosso cérebro preenche o buraco, revestindo-o, tirando uma média. Isso significa que parte da nossa visão é falsa, gerada pela mente subconsciente para nos enganar.

E também só enxergamos com clareza o centro do nosso campo de visão, chamado de fóvea. A parte periférica é embaçada para economizar energia. Mas a fóvea é muito pequena. Para captar a maior quantidade de informação com a minúscula fóvea, os olhos se movem o tempo todo. Esse movimento de nossos olhos, rápido e inquieto, é chamado de sacada ou movimento sacádico. É produzido no subconsciente, dando-nos a falsa impressão de que nosso campo de visão é claro e bem focado.

Quando menino, a primeira vez que vi um diagrama do espectro eletromagnético, fiquei chocado. Eu não tinha a menor ideia de que grande parte de tal espectro (por exemplo, luz infravermelha e ultravioleta, raios X, raios gama) era totalmente invisível para nós. Comecei a perceber que o que estava vendo com meus próprios olhos era apenas uma aproximação minúscula e grosseira da realidade. (Segundo um velho ditado, “Se aparência fosse o mesmo que essência, não precisaríamos da ciência”). Temos sensores na retina que só conseguem detectar vermelho, verde e azul. Isso significa que nunca vimos realmente o amarelo, o marrom, o laranja, e uma série de outras cores. Essas cores existem de fato, mas nosso cérebro só consegue fazer uma aproximação de cada uma delas, misturando diferentes quantidades de vermelho, verde e azul. (É possível constatar isso olhando muito atentamente para uma tela de TV em cores antiga. Vemos apenas um monte de pontinhos vermelhos, verdes e azuis. A TV em cores é de fato uma ilusão.)

Nossos olhos enganam também nos levando a pensar que podemos ver a profundidade. As retinas dos olhos são bidimensionais, mas como temos dois olhos separados por alguns centímetros, os cérebros direito e esquerdo fundem as duas imagens, dando-nos a falsa sensação de uma terceira dimensão. Para objetos mais distantes, podemos julgar quão distante está o objeto observando seu movimento conforme movemos a cabeça. Isso se chama paralaxe.

Essa paralaxe explica o fato de crianças acharem que “a lua as está seguindo”. Como o cérebro tem dificuldade em compreender a paralaxe de um objeto tão distante quanto a Lua, parece que a Lua está sempre a uma distância fixa “atrás” de nós, mas é apenas uma ilusão causada pelo cérebro tomando um atalho.

O PARADOXO DO CÉREBRO DIVIDIDO

Uma diferença entre essa imagem – baseada na hierarquia corporativa de uma empresa – e a verdadeira estrutura do cérebro pode ser observada no caso curioso de pacientes com cérebro dividido. Um traço incomum do cérebro é que ele tem duas metades, ou hemisférios, o direito e o esquerdo, praticamente idênticos. Durante muito tempo os cientistas se perguntaram por que o cérebro tem essa redundância desnecessária, pois consegue funcionar mesmo quando um hemisfério é totalmente removido. Nenhuma hierarquia corporativa apresenta essa estranha característica. Além disso, se cada hemisfério tem consciência, isso significa que temos dois centros de consciência dentro do mesmo crânio?

O dr. Roger W. Sperry, do California Institute of Technology, ganhou o Prêmio Nobel em 1981 ao mostrar que os dois hemisférios cerebrais não são cópias exatas um do outro, e de fato desempenham funções diferentes. Essa descoberta causou furor na neurologia (e gerou uma indústria duvidosa de livros de autoajuda que pregam a aplicação da dicotomia cérebro-esquerdo/cérebro-direito na vida diária.)

O dr. Sperry estava tratando de epiléticos, que às vezes sofrem de convulsões do tipo “grande mal” provocadas pelo descontrole do processo contínuo de realimentação (ou ciclos de feedback) entre os dois hemisférios do cérebro. Tais convulsões – causadas pelo mesmo processo que gera o som estridente da microfonia – podem até matar. O dr. Sperry começou por cortar o corpo caloso, que liga os dois hemisférios do cérebro, de modo que não mais se comunicassem, e assim não compartilhassem informações entre os lados direito e esquerdo do corpo. Isso geralmente fazia parar o processo de realimentação e as convulsões.

A princípio, os pacientes com cérebro dividido pareciam perfeitamente normais. Continuavam alertas e mantinham uma conversa naturalmente, como

se nada tivesse acontecido. Porém, uma análise mais minuciosa desses indivíduos mostrou que havia algo muito diferente neles.

Normalmente, os hemisférios se complementam, com os pensamentos indo e vindo de um a outro. O cérebro esquerdo é mais analítico e lógico. É onde se encontram as habilidades verbais, ao passo que o direito é mais holístico e artístico. Mas o cérebro esquerdo é dominante, é o que toma as decisões finais. Os comandos passam do cérebro esquerdo para o direito por meio do corpo caloso. Se essa conexão é cortada, o cérebro direito fica livre da ditadura do esquerdo. Talvez o cérebro direito tenha vontade própria, contrariando os desejos do esquerdo dominante.

Resumindo, pode haver duas vontades agindo dentro do mesmo crânio, às vezes brigando pelo controle do corpo. Isso cria a estranha situação em que a mão esquerda (controlada pelo hemisfério direito) começa a agir de forma independente de nossos desejos, como se fosse um apêndice externo.

Há um caso documentado de um homem que estava a ponto de abraçar sua esposa com uma das mãos, quando descobriu que a outra mão tinha uma intenção diferente: dar um soco no rosto dela. Uma mulher relatou que estava pegando um vestido com uma das mãos enquanto a outra mão pegava uma roupa totalmente diferente. E um homem não conseguia dormir à noite com medo de que a mão rebelde fosse estrangulá-lo.

Às vezes, pessoas com cérebro dividido pensam que estão vivendo em um desenho animado, com uma das mãos tentando controlar a outra. Alguns médicos chamam isso de “síndrome do dr. Fantástico” [dr. Strangelove], por causa da cena do filme em que uma das mãos do doutor luta contra a outra.

Após estudos detalhados de pacientes com cérebro dividido, o dr. Sperry concluiu que poderia haver duas mentes distintas operando num único cérebro. Ele escreveu que cada hemisfério é “de fato um sistema consciente em si mesmo, capaz de perceber, pensar, lembrar, raciocinar, querer, se emocionar, tudo isso num nível caracteristicamente humano, e (...) os dois hemisférios podem estar passando por experiências mentais diferentes, e até conflitantes, ao mesmo tempo”.

Quando entrevistei o dr. Michael Gazzaniga, da Universidade da Califórnia em Santa Bárbara, uma autoridade em pacientes de cérebro dividido, perguntei como podem ser feitos experimentos para testar essa teoria. Há várias formas de comunicação com cada hemisfério sem o conhecimento do outro. Pode-se, por exemplo, colocar no sujeito óculos especiais em que aparecem perguntas diferentes diante de cada olho, de modo que é fácil direcionar perguntas a cada hemisfério separadamente. O difícil é tentar obter uma resposta de cada hemisfério. Como o cérebro direito não consegue falar (os centros da fala estão situados no lado esquerdo) é difícil obter suas respostas. Dr. Gazzaniga me disse que, para descobrir o que o cérebro direito estava pensando, ele criou um

experimento em que esse hemisfério (mudo) conseguia “falar” usando peças com letras, como num jogo de palavras cruzadas.

Ele começou perguntando ao cérebro esquerdo do paciente o que ele iria fazer depois que se formasse. O paciente respondeu que queria ser projetista. Tudo ficou mais interessante quando a mesma pergunta foi feita ao cérebro direito (mudo), que soletrou as palavras “piloto de carro de corrida”. Sem o conhecimento do cérebro esquerdo dominante, o cérebro direito tinha planos totalmente diferentes para o futuro. O cérebro direito tinha, literalmente, uma mente própria.

Rita Carter escreve: “As implicações possíveis nos deixam aturdidos. Sugerem que todos nós podemos estar carregando por aí um prisioneiro mudo dentro do crânio, com uma personalidade, ambições e consciência de si muito diferentes das que acreditamos ter no cotidiano.”

Talvez haja verdade na afirmação de que “dentro dele existe alguém ansiando por ser livre”. Isso significa que os dois hemisférios podem ter crenças diferentes. Por exemplo: o neurologista V. S. Ramachandran descreve um paciente de cérebro dividido que, ao ser perguntado se era religioso, disse que era ateu, mas o cérebro direito declarou o contrário. Pelo visto, é possível ter duas posições religiosas opostas residindo no mesmo cérebro. Ramachandran prossegue: “O que acontece quando essa pessoa morre? Um hemisfério vai para o céu e o outro vai para o inferno? Não sei a resposta.”

É concebível, portanto, que uma pessoa de cérebro dividido possa ser republicana e democrata ao mesmo tempo. Se perguntarmos em quem ela vai votar, ela dirá o candidato do cérebro esquerdo, pois o direito não consegue falar. Mas dá para imaginar o caos na cabine de votação quando essa pessoa tem que usar só uma das mãos.

QUEM ESTÁ NO COMANDO?

Um dos que dedicou seu tempo e fez muitas pesquisas para entender o problema da mente subconsciente é o dr. David Eagleman, neurocientista do Baylor College of Medicine. Quando o entrevistei, perguntei: “Se a maior parte dos nossos processos mentais são subconscientes, por que ignoramos esse fato tão importante?” Ele deu o exemplo de um jovem rei que herda o trono e leva o crédito de tudo o que acontece no reino, mas não tem a menor noção da enorme quantidade de assessores, camponeses e soldados necessários para manter o trono.

Nossas escolhas de políticos, cônjuges, amigos e profissões são influenciadas por coisas de que não temos consciência. (Por exemplo: é um resultado muito

estranho, ele diz, que “pessoas chamadas Denise ou Denis tenham uma probabilidade muito maior de se tornarem dentistas; que pessoas chamadas Laura ou Lawrence tenham maior probabilidade de serem advogadas [*lawyers*] e pessoas chamadas George ou Georgiana, de se tornarem geólogas”). Isso significa também que o que consideramos ser a “realidade” é apenas uma aproximação feita pelo cérebro para preencher as lacunas. Cada um de nós vê a realidade de maneira ligeiramente diferente. Por exemplo, ele observa, “pelo menos 15% das mulheres possuem uma mutação genética que lhes dá um tipo extra (um quarto tipo) de fotorreceptor de cor – e isso lhes permite distinguir entre cores que parecem idênticas para a maioria das pessoas, que têm apenas três tipos de fotorreceptor.”

Certamente, quanto mais entendemos a mecânica do pensamento, mais perguntas aparecem. O que acontece exatamente no centro de comando da mente quando confrontado com um obscuro centro de comando rebelde? Afinal, o que queremos dizer com “consciência”, se ela pode ser dividida pela metade? E qual é a relação entre consciência e “si mesmo”, e “consciência de si”?

Se pudermos responder a essas difíceis questões, talvez a resposta possa nos indicar o caminho para o entendimento de consciência não humana, a consciência de robôs e alienígenas do espaço sideral, por exemplo, que pode ser inteiramente diferente da nossa.

Vamos então propor uma resposta clara para essa enganadora e complexa questão: O que é consciência?

A mente do homem é capaz de tudo... porque tudo está nela, todo o passado e também todo o futuro.

– JOSEPH CONRAD

A consciência pode reduzir até o mais cuidadoso pensador a um falastrão incoerente.

– COLIN MCGINN

2 CONSCIÊNCIA – O PONTO DE VISTA DE UM FÍSICO

A ideia de consciência tem intrigado filósofos há séculos, mas ela tem resistido a uma definição simples até os dias de hoje. O filósofo David Chalmers catalogou mais de vinte mil artigos sobre o tema. Nunca na ciência tantos se dedicaram tanto para criar tão pouco consenso. O pensador do século XVII Gottfried Leibniz escreveu certa vez: “Se fosse possível ampliar o cérebro até ficar do tamanho de um moinho e andar lá por dentro, não encontraríamos a consciência.”

Alguns filósofos duvidam de que seja possível uma teoria da consciência. Alegam que a consciência não pode ser explicada, dado que um objeto não pode entender a si mesmo; portanto, não temos sequer o poder mental para resolver essa questão desconcertante. O psicólogo Steven Pinker, de Harvard, escreveu: “Não conseguimos ver a luz ultravioleta. Não conseguimos girar mentalmente um objeto na quarta dimensão. E talvez não consigamos entender enigmas como o livre-arbítrio e a consciência.”

De fato, na maior parte do século XX, uma das teorias dominantes da psicologia, o behaviorismo, negou totalmente a importância da consciência. O behaviorismo se baseia na ideia de que somente o comportamento objetivo de animais e pessoas é digno de estudo, e não os estados internos, subjetivos da mente.

Outros desistiram de definir a consciência e tentam simplesmente descrevê-la. O psiquiatra Giulio Tononi disse: “Todo mundo sabe o que é consciência: é aquilo que nos abandona toda noite quando dormimos sem sonhar e retorna na manhã seguinte quando acordamos.”

Embora a natureza da consciência esteja sendo discutida há séculos, pouco se concluiu. Já que muitas das invenções que possibilitaram os avanços da ciência do cérebro foram criadas por físicos, talvez seja útil usar um exemplo da física para reexaminar essa velha questão.

COMO OS FÍSICOS ENTENDEM O UNIVERSO

Quando um físico tenta entender alguma coisa, primeiro reúne dados e depois propõe um “modelo”, uma versão simplificada do objeto de estudo, que capta suas características essenciais. Na física, o modelo é descrito por uma série de parâmetros (por exemplo, temperatura, energia, tempo). Depois o físico usa esse modelo para prever sua evolução, simulando seus movimentos. De fato, alguns dos maiores supercomputadores do mundo são usados para simular a evolução de modelos, que podem descrever prótons, explosões nucleares, padrões climáticos, o Big Bang e o centro de buracos negros. Depois criam um modelo

melhor, usando parâmetros mais sofisticados e continuam fazendo simulações.

Por exemplo: quando Isaac Newton estava analisando o movimento da Lua, criou um modelo simples, mas que iria mudar o curso da história da humanidade. Ele imaginou lançar uma maçã no ar. Quanto mais rápido a maçã fosse atirada, ele raciocinou, mais longe iria cair. Se atirasse com rapidez suficiente, ela poderia dar a volta ao mundo e até retornar ao ponto inicial. Depois Newton afirmou que esse modelo representava a trajetória da Lua, e as forças que guiavam o movimento da maçã em torno da Terra eram idênticas às forças guiando a Lua.

Mas esse modelo sozinho ainda era inútil. O ponto de virada foi quando Newton conseguiu usar essa teoria para simular o futuro, para calcular a posição futura de objetos em movimento. Era um problema difícil, exigindo que ele criasse um campo inteiramente novo da matemática, chamado cálculo. Usando essa nova matemática, Newton conseguiu prever a trajetória não só da Lua, mas também do cometa Halley e dos planetas. Desde então os cientistas têm usado as leis de Newton para simular a futura trajetória de objetos em movimento, sejam balas de canhão, máquinas, automóveis, foguetes, asteroides, meteoros, ou até estrelas e galáxias.

O sucesso ou fracasso de um modelo depende de quão fielmente ele reproduz os parâmetros básicos do original. Nesse caso o parâmetro básico era a localização da maçã e da Lua no espaço e no tempo. Ao deixar este parâmetro evoluir (isto é, ao longo do tempo), Newton desvendou, pela primeira vez na história, a ação de corpos em movimento, que é uma das mais importantes descobertas na ciência.

Os modelos são úteis até serem substituídos por modelos mais precisos, descritos por parâmetros melhores. Einstein substituiu a imagem das forças de Newton agindo sobre maçãs e luas por um novo modelo baseado em um novo parâmetro, a curvatura do tempo e do espaço. A maçã não se movia porque a Terra exercia uma força sobre ela, mas porque o tecido do espaço e do tempo era esticado pela Terra, de modo que a maçã estava simplesmente se movendo ao longo da superfície de um espaço-tempo curvo. A partir disso, Einstein pôde simular o futuro de todo o universo. Agora, com computadores, podemos fazer simulações desse modelo no futuro e criar belas imagens de colisões de buracos negros.

Vamos incorporar essa estratégia básica a uma nova teoria da consciência.

DEFINIÇÃO DE CONSCIÊNCIA

Recolhi descrições dos campos da neurologia e da biologia para definir

consciência da seguinte maneira:

Consciência é o processo de criar um modelo do mundo usando múltiplos ciclos de feedback em vários parâmetros (por exemplo, temperatura, espaço, tempo e em relação a outros), a fim de atingir uma meta (por exemplo, encontrar parceiros, comida, abrigo).

Chamo isso de “teoria espaço-tempo da consciência”, porque ela enfatiza a ideia de que animais criam um modelo do mundo, principalmente em relação ao espaço e entre outros animais, enquanto os humanos vão além e criam um modelo do mundo em relação ao tempo, tanto futuro como passado.

Por exemplo: o nível mais baixo de consciência é o Nível 0, em que um organismo é estacionário, ou tem mobilidade limitada, e cria um modelo de seu lugar usando o ciclo de feedback de alguns poucos parâmetros (por exemplo, temperatura). Nesse exemplo, o nível mais simples de consciência é um termostato. Ele liga automaticamente o ar-condicionado ou o aquecedor para ajustar a temperatura do ambiente, sem ajuda. A chave é um ciclo de feedback que liga um interruptor se a temperatura ficar muito quente ou muito fria. (Por exemplo: metais se expandem quando aquecidos, e assim o termostato pode acionar um interruptor se uma tira de metal se expande além de certo ponto.)

Cada ciclo de feedback registra “uma unidade de consciência”, portanto um termostato teria uma só unidade de consciência de Nível 0, ou seja, Nível 0:1.

Desse modo, podemos classificar a consciência numericamente, com base no número e na complexidade dos ciclos de feedback usados para criar um modelo do mundo. A consciência então já não é uma vaga compilação de conceitos indefinidos, circulares, mas um sistema de hierarquias que podem ser classificadas numericamente. Uma bactéria ou uma flor, por exemplo, têm muito mais ciclos de feedback, portanto têm um grau mais alto dentro do Nível 0 de consciência. Uma flor com dez ciclos de feedback (que medem a temperatura, umidade, luz solar, gravidade etc.), tem um Nível 0:10 de consciência.

Organismos com mobilidade e um sistema nervoso central têm Nível I de consciência, que inclui um novo conjunto de parâmetros para medir suas mudanças de lugar. Um exemplo de Nível I de consciência são os répteis. Eles têm tantos ciclos de feedback que desenvolveram um sistema nervoso central para lidar com isso. O cérebro reptiliano tem talvez cem ou mais ciclos de feedback (controlando o sentido do olfato, equilíbrio, tato, audição, visão, pressão sanguínea etc., e cada um desses contém mais ciclos de feedback). Por exemplo: a visão envolve um grande número de ciclos de feedback, pois os olhos reconhecem cores, movimentos, formas, intensidade de luz e sombras. Da mesma forma, os outros sentidos dos répteis, como a audição e o paladar,

exigem ciclos de feedback adicionais. A totalidade desses numerosos ciclos de feedback cria uma imagem mental de onde o réptil está situado no mundo e onde outros animais (por exemplo, a presa) estão situados. Assim, o Nível I de consciência é controlado principalmente pelo cérebro reptiliano, situado no fundo e no centro do cérebro humano.

A seguir, temos o Nível II de consciência, em que os organismos criam um modelo de seu lugar não só no espaço, mas também em relação a outros (isto é, são animais sociais com emoções). A quantidade de ciclos de feedback do Nível II de consciência aumenta exponencialmente, portanto convém introduzir uma outra classificação numérica para esse tipo de consciência. Adquirir aliados, detectar inimigos, obedecer ao macho alfa etc. são comportamentos muito complexos, que exigem um cérebro muito expandido, portanto o Nível II de consciência coincide com a formação de novas estruturas cerebrais, na forma do sistema límbico. Como vimos, o sistema límbico inclui o hipocampo (para a memória), a amígdala (para as emoções) e o tálamo (para a informação sensorial), e todos fornecem novos parâmetros para criar modelos em relação a outros. Portanto, o número e os tipos de ciclos de feedback mudam.

Definimos o grau do Nível II de consciência como o número total de diferentes ciclos de feedback necessários para um animal interagir socialmente com membros do seu grupo. Infelizmente, estudos de consciência animal são extremamente limitados, por isso muito pouco se fez para catalogar todos os meios pelos quais os animais se comunicam uns com os outros. Mas de um modo geral, podemos estimar o Nível II de consciência contando o número de animais num bando ou grupo e listando todas as formas de interação emocional entre eles. Isso inclui reconhecer amigos e rivais, formar laços com outros, retribuir favores, formar coalizões, entender o próprio status e a posição social de outros, respeitar os superiores, demonstrar poder sobre os inferiores, conspirar para subir os degraus sociais etc. (Excluimos insetos do Nível II porque, embora tenham relações sociais com os membros do grupo ou da colmeia, até onde sabemos não têm emoções.)

Apesar da falta de estudos empíricos de comportamento animal, podemos chegar a uma classificação numérica, muito grosseira, do Nível II de consciência, listando todas as emoções e comportamentos sociais demonstrados. Por exemplo: se uma alcateia tem dez lobos e cada um deles interage com todos os outros por meio de quinze emoções e gestos diferentes, então o nível de consciência do lobo, num primeiro cálculo, é dado pelo produto dos dois, isto é, 150, o que dará Nível II:150 de consciência. Esse valor leva em conta o número de animais com que ele interage e o número de formas de comunicação com cada um dos outros. Esse valor é apenas uma aproximação do total de interações sociais que o animal apresenta, e certamente será mudado à medida que aprendermos mais sobre seu comportamento.

É claro que, como a evolução nunca é precisa nem nítida, há casos especiais que precisamos explicar, como o nível de consciência de animais sociais que são caçadores solitários. Explicaremos isso nas Notas.

NÍVEL III DE CONSCIÊNCIA: SIMULANDO O FUTURO

Dada essa estrutura para a consciência, vemos que os humanos não são únicos, e que há um *continuum* de consciência. Como Charles Darwin comentou certa vez: “A diferença entre o homem e os animais mais desenvolvidos, por maior que seja, é certamente de grau, e não de tipo.” Mas o que separa a consciência humana da consciência dos animais? Os homens são os únicos no reino animal que entendem o conceito de amanhã. Ao contrário dos animais, sempre nos perguntamos “E se?” em semanas, meses ou daqui a anos. Portanto, acredito que o Nível III de consciência cria um modelo de seu lugar no mundo e então simula esse modelo no futuro fazendo previsões aproximadas. Podemos resumir isso da seguinte maneira:

A consciência humana é uma forma específica de consciência, que cria um modelo do mundo e o simula ao longo do tempo, avaliando o passado para simular o futuro. Isso exige mediar e avaliar muitos ciclos de feedback a fim de tomar uma decisão para atingir uma meta.

Quando atingimos o Nível III de consciência, são tantos os ciclos de feedback que precisamos de um diretor-geral para analisá-los a fim de simular o futuro e tomar uma decisão final. Por isso, nosso cérebro difere dos cérebros dos animais, principalmente na expansão do córtex pré-frontal, situado logo atrás da testa, que nos permite “ver” o futuro.

O dr. Daniel Gilbert, psicólogo de Harvard, escreveu: “A maior conquista do cérebro humano é a capacidade de imaginar objetos e episódios que não existem no reino do real, e é essa capacidade que nos permite pensar no futuro. Como um filósofo observou, o cérebro humano é uma ‘máquina de antecipação’, e ‘fazer o futuro’ é a sua função mais importante.”

Usando varreduras cerebrais, podemos até supor qual é a área exata do cérebro onde ocorre a simulação do futuro. O neurologista Michael Gazzaniga observa que a “área 10 (a camada granular interna IV), no córtex pré-frontal lateral, é quase duas vezes maior nos humanos do que nos macacos. A área 10 está ligada à memória e ao planejamento, à flexibilidade cognitiva, ao pensamento abstrato, ao acionamento do comportamento adequado e à inibição do comportamento inadequado, ao aprendizado de regras e à seleção das

informações relevantes do que é percebido por meio dos sentidos”. (Neste livro, vamos nos referir a essa área, onde está concentrada a tomada de decisões, como córtex pré-frontal dorsolateral, apesar de ter alguma sobreposição com outras áreas do cérebro.)

Embora os animais possam ter um entendimento bem definido de seu lugar no espaço, e alguns tenham certo grau de consciência de outros animais, não está claro se planejam sistematicamente o futuro nem se têm algum entendimento do “amanhã”. Muitos animais, inclusive animais sociais com sistema límbico bem desenvolvido, reagem a situações (por exemplo, a presença de predadores ou de parceiros em potencial) confiando principalmente no instinto, e não planejando sistematicamente o futuro.

Por exemplo: os mamíferos não se preparam para o inverno planejando a hibernação, apenas seguem instintivamente as quedas de temperatura. Um ciclo de feedback regula a hibernação. A consciência é dominada por mensagens vindas dos sentidos. Não há evidências de que analisem sistematicamente diversos planos e arranjos quando se preparam para hibernar. Os predadores, quando usam a astúcia e o disfarce para surpreender uma presa incauta, antecipam os eventos futuros, mas esse planejamento se limita ao instinto e à duração da caçada. Os primatas elaboram planos de curto prazo (por exemplo, para encontrar comida), mas não há indicação de que planejem isso com mais que algumas horas de antecedência.

Os humanos são diferentes. Apesar de confiarmos no instinto e nas emoções em muitas situações, também analisamos e avaliamos informações de vários ciclos de feedback. Fazemos isso organizando simulações, às vezes até além do nosso tempo de vida, e mesmo milhares de anos no futuro. A vantagem de fazer simulações é avaliar as várias possibilidades, a fim de tomar a melhor decisão para atingir um objetivo. Isso ocorre no córtex pré-frontal, que nos permite simular o futuro e avaliar as possibilidades para traçar a melhor linha de ação.

Essa capacidade evoluiu por várias razões. Primeiro, ter a capacidade de considerar o futuro traz enormes benefícios evolucionários, como fugir de predadores, encontrar comida e parceiros. Segundo, nos permite escolher entre vários resultados diferentes e selecionar o melhor.

Terceiro, o número de ciclos de feedback aumenta exponencialmente à medida que passamos do Nível 0 para o Nível I e para o Nível II, portanto precisamos de um “diretor-geral” para avaliar todas as mensagens conflitantes. O instinto já não é suficiente. É preciso haver um corpo central para avaliar cada um dos ciclos de feedback. Isso diferencia a consciência humana da dos animais. Esses ciclos de feedback, por sua vez, são avaliados com simulações do futuro para se obter o melhor resultado. Se não tivéssemos um diretor, reinaria o caos e teríamos uma sobrecarga sensorial.

Um experimento simples pode demonstrar isso. David Eagleman observa que

se pegamos um peixe esgana-gata macho e colocamos uma fêmea cruzando seu território, o macho fica confuso porque quer se acasalar com a fêmea, mas ao mesmo tempo quer defender seu território. Sendo assim, o esgana-gata ataca a fêmea enquanto tenta fazer-lhe a corte. O macho fica enlouquecido, tentando cortejar e ao mesmo tempo matar a fêmea.

Isso acontece também com os camundongos. Coloque um eletrodo na frente de um pedaço de queijo. Se o ratinho chegar muito perto, vai levar um choque. Um ciclo de feedback diz a ele para comer o queijo, mas outro diz que se afaste para evitar o choque. Ajustando a localização do eletrodo, você pode fazer o ratinho oscilar, dividido entre dois ciclos de feedback conflitantes. Enquanto o homem tem um diretor-geral no cérebro para avaliar os prós e contras da situação, o rato, guiado por dois ciclos de feedback conflitantes, vai para frente e para trás. (É como o provérbio do burro que morre de fome porque foi colocado entre dois montes iguais de feno.)

Como, exatamente, o cérebro simula o futuro? O cérebro humano é inundado por uma grande quantidade de dados sensoriais e emocionais. A chave é simular o futuro fazendo conexões causais entre eventos – isto é, se acontecer A, então acontece B. Mas se acontecer B, pode resultar em C e D. Isso dispara uma reação em cadeia de eventos, acabando por criar uma árvore de futuros possíveis encadeados, com muitas ramificações. O diretor no córtex pré-frontal avalia os resultados dessas árvores causais a fim de tomar a decisão final.

Digamos que você queira assaltar um banco. Quantas simulações realistas desse evento é possível fazer? É preciso pensar nas várias conexões causais envolvendo a polícia, clientes, sistemas de alarme, relações com os comparsas, condições do trânsito, a proximidade da delegacia etc. Para uma simulação bem-sucedida do assalto, seria necessário analisar centenas de conexões causais.

É também possível medir numericamente esse nível de consciência. Digamos que se apresente a uma pessoa uma série de situações diferentes, como a descrita acima, para ela simular o futuro de cada uma. A soma total das conexões causais em que tal pessoa consegue pensar para todas essas situações pode ser tabulada. Uma dificuldade é que há um número ilimitado de conexões causais que poderiam ser feitas para cada variedade de situação. Para contornar tal dificuldade, dividimos esse número pela média de conexões causais obtida de um grande grupo de controle. Como num teste de Q.I., podemos multiplicar esse número por 100. Assim, o nível de consciência de alguém, por exemplo, pode ser Nível III:100, o que significa que é capaz de simular o futuro como uma pessoa comum.

Resumimos esses níveis de consciência no seguinte diagrama:

NÍVEIS DE CONSCIÊNCIA DE DIFERENTES ESPÉCIES



Nível	Espécie	Parâmetro	E
0	Plantas	Temperatura, luz solar	N
I	Répteis	Espaço	T ce
II	Mamíferos	Relações sociais	Si lí:
III	Humanos	Tempo (esp. futuro)	C p1 fr

Teoria espaço-tempo da consciência. Definimos consciência como o processo de criar um modelo do mundo usando múltiplos ciclos de feedback de vários parâmetros (por exemplo, no espaço, no tempo, e em relação aos outros), a fim de atingir um objetivo. A consciência humana é um tipo particular que envolve a mediação entre esses ciclos de feedback por meio de simulação do futuro e avaliação do passado.

Observem que essas categorias correspondem aproximadamente aos níveis evolucionários que encontramos na natureza – por exemplo, répteis, mamíferos e humanos. Entretanto, há também áreas cinzentas, como animais que possuem alguns aspectos de diferentes níveis de consciência, animais que fazem planejamentos rudimentares, ou mesmo células que se comunicam entre si. Esse

diagrama pretende apenas dar uma visão maior, global, de como a consciência está organizada no reino animal.

O QUE É HUMOR? POR QUE TEMOS EMOÇÕES?

Toda teoria tem que ser contestável. O desafio para a teoria do espaço-tempo da consciência é explicar *todos* os aspectos da consciência humana nesse esquema. Pode ser contestável se houver padrões de pensamento que não podem ser encaixados na teoria. Um crítico diria que certamente nosso senso de humor é tão esquisito e efêmero que está além de qualquer explicação. Passamos muito tempo rindo com amigos e comediantes; entretanto parece que o humor nada tem a ver com simulações do futuro. Mas considerem que muito do humor, como contar uma piada, depende da sacada final.

Ao ouvir uma piada, não deixamos de simular o futuro e completar a história (mesmo quando fazemos isso sem perceber). Conhecemos o suficiente do mundo físico e social para antecipar o final, por isso damos uma gargalhada quando o final da piada tem uma conclusão totalmente inesperada. A essência do humor é quando nossa simulação do futuro é interrompida de repente, de maneira surpreendente. Isso teve importância histórica para nossa evolução porque o sucesso depende, em parte, de nossa capacidade de simular eventos futuros. Como a vida na selva é cheia de eventos não previstos, quem é capaz de prever resultados inesperados tem maior chance de sobreviver. Assim, ter um senso de humor bem desenvolvido é de fato uma indicação de inteligência e consciência de Nível III, isto é, a capacidade de simular o futuro.

Por exemplo: Uma vez pediram a W. C. Fields uma opinião sobre atividades sociais para jovens. A pergunta foi “Você acredita em *clubs* para jovens?”, e ele respondeu “Só quando falha a gentileza” ^[1] A piada tem graça porque simulamos mentalmente um futuro com crianças num clube social, enquanto W. C. Fields simulou um futuro diferente, envolvendo uma arma. (É claro que uma piada explicada perde o poder, pois já simulamos mentalmente vários futuros possíveis.)

Isso mostra também o que todo comediante sabe: o momento certo é o segredo do humor. Se o final é contado muito depressa, o cérebro não teve tempo de simular o futuro, e a experiência do imprevisto não acontece. Se demorar demais, o cérebro tem tempo de simular vários futuros possíveis, e o final perde o elemento surpresa.

O humor tem outras funções, é claro, como criar laços com os membros do grupo. De fato, usamos o senso de humor como um meio de formar opinião sobre o caráter dos outros. Isso, por sua vez, é essencial para determinar nosso

status numa sociedade. Portanto, a gargalhada também ajuda a definir nossa posição no mundo social, isto é, no Nível II de consciência.

POR QUE BRINCAMOS E FOFOCAMOS?

Até atividades aparentemente triviais, como fofocar ou fazer brincadeiras com amigos, devem ser explicadas nesse quadro teórico. (Se um marciano visse a enorme oferta de revistas de fofoca perto do caixa em um supermercado, poderia concluir que fofocar é a principal atividade dos humanos, observação esta que não estaria longe da verdade.)

A fofoca é essencial para a sobrevivência porque, como os complexos mecanismos das interações sociais estão sempre mudando, ela é uma forma de entender esse campo social em constante mutação. Isso é o Nível II de consciência em ação. Mas quando ouvimos um mexerico, fazemos simulações imediatas para saber como isso pode afetar nossa própria posição na sociedade, o que nos leva ao Nível III de consciência. Na verdade, milhares de anos atrás, a fofoca era a única maneira de obter informações vitais sobre a tribo. A própria vida da pessoa dependia de saber da última fofoca.

Algo tão superficial quanto “brincar” é também um aspecto essencial da consciência. Quando perguntamos a uma criança por que ela gosta de brincar, ela responde “Porque é divertido”. Mas isso leva à pergunta seguinte: O que é diversão? Na verdade, quando as crianças brincam, elas geralmente recriam, de forma simplificada, interações humanas complexas. A sociedade humana é extremamente sofisticada, complicada demais para o cérebro infantil ainda em desenvolvimento, então as crianças fazem simulações simplificadas da sociedade dos adultos, brincando de médico, polícia e ladrão, escolinha. Cada brincadeira é um modelo que permite à criança vivenciar um pequeno segmento do comportamento adulto e fazer simulações do futuro. Da mesma forma, quando adultos participam de jogos, como o pôquer, o cérebro cria continuamente um modelo de quais cartas os outros jogadores possuem e projeta esse modelo no futuro, usando dados anteriores sobre a personalidade dos parceiros de jogo, a tendência a blefar etc. O segredo de jogos como o xadrez, baralho e jogos de azar é a capacidade de simular o futuro. Os animais, que vivem muito no presente, não são tão bons nos jogos quanto os humanos, principalmente nos jogos que envolvem planejamento. Filhotes de mamíferos praticam formas de jogos, mas é mais como exercício, testando uns aos outros, praticando para lutas futuras e estabelecendo uma hierarquia social, e não simulando o futuro.

Minha teoria do espaço-tempo da consciência pode também ajudar a esclarecer outro tópico controverso: a inteligência. Embora os testes de Q.I.

afirmem medir a “inteligência”, na verdade não oferecem sequer uma definição inicial de inteligência. De fato, um cínico diria, com alguma razão, que o Q.I. é uma medida de “quanto alguém se sai bem em testes de Q.I.”, o que é circular. Além disso, os testes de Q.I. são criticados pelo excesso de referências culturais. Porém, nessa nova perspectiva, a inteligência pode ser vista como a complexidade de nossas simulações do futuro. Consequentemente, um mestre do crime, que pode ter abandonado a escola, ser um analfabeto funcional e obter uma pontuação baixíssima num teste de Q.I., pode muito bem superar a capacidade da polícia. Ser mais esperto que a polícia pode implicar simplesmente fazer simulações mais sofisticadas do futuro.

NÍVEL I: FLUXO DE CONSCIÊNCIA

Os humanos são provavelmente os únicos no planeta capazes de operar em todos os níveis de consciência. Usando IRM, podemos analisar as diferentes estruturas envolvidas em cada nível de consciência.

Para nós, o fluxo do Nível I de consciência é basicamente a interação entre o córtex pré-frontal e o tálamo. Quando passeamos tranquilamente pelo parque, percebemos o cheiro das plantas, a brisa suave, os estímulos visuais do sol etc. Nossos sentidos enviam sinais para a medula espinhal, o tronco cerebral e então ao tálamo, que opera como uma estação de transmissão, classificando os estímulos e os enviando para os vários córtices cerebrais. As imagens do parque, por exemplo, são enviadas para o córtex occipital, na parte de trás do cérebro, e a sensação da brisa na pele é enviada para o lobo parietal. Os sinais são processados nos córtices apropriados e então enviados para o córtex pré-frontal, onde finalmente nos tornamos conscientes de todas essas sensações.

Isso é ilustrado na Figura 7.

NÍVEL II: EM BUSCA DE UM LUGAR NA SOCIEDADE

Enquanto o Nível I de consciência usa as sensações para criar um modelo de nossa localização física no espaço, o Nível II de consciência cria um modelo de nosso lugar na sociedade.

Imagine que vamos a uma festa elegante, em que as pessoas essenciais para nosso trabalho estão presentes. Passando os olhos pela sala, tentando identificar os colegas, há uma interação intensa entre o hipocampo (que processa a memória), a amígdala (que processa as emoções) e o córtex pré-frontal (que junta as

informações).

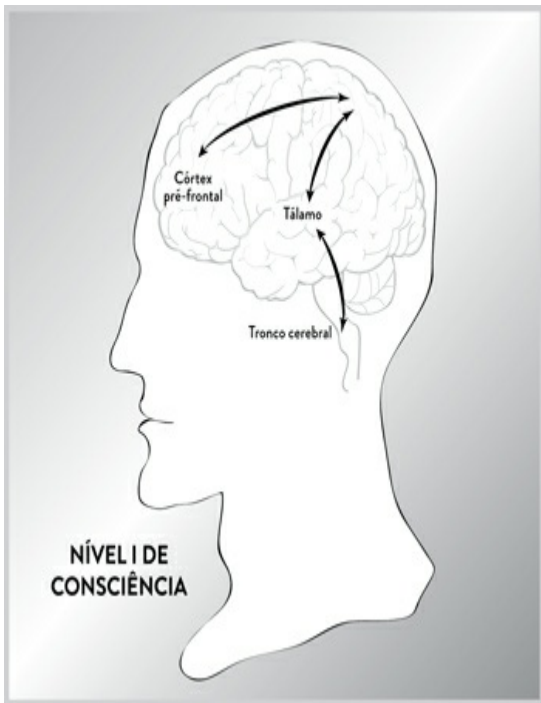


Figura 7. No Nível I de consciência, a informação sensorial viaja pelo tronco cerebral, passa pelo tálamo, segue pelos vários córtices do cérebro e finalmente chega ao córtex pré-frontal. Assim, o fluxo do Nível I de consciência é criado pelo fluxo de informação do tálamo para o córtex pré-frontal.

A cada imagem o cérebro anexa automaticamente uma emoção, como felicidade, medo, raiva, ciúme, e processa a emoção na amígdala.

Se avistamos nosso maior rival na festa, aquele que suspeitamos nos apunhalar pelas costas, a emoção do medo é processada pela amígdala, que envia uma mensagem urgente para o córtex pré-frontal, alertando do possível perigo. Ao mesmo tempo, são enviados sinais ao sistema endócrino para começar a bombear adrenalina e outros hormônios para o sangue, aumentando assim os batimentos cardíacos e nos preparando para uma possível reação de luta ou fuga.

Isso é ilustrado na Figura 8.

Mas além de simplesmente reconhecer pessoas, o cérebro tem a fantástica capacidade de imaginar o que os outros estão pensando. Isso é chamado de Teoria da Mente, proposta inicialmente pelo dr. David Premack, da Universidade da Pensilvânia, que é a capacidade de deduzir os pensamentos dos outros. Em toda sociedade complexa, quem tem a capacidade de adivinhar corretamente intenções, motivos e planos de outras pessoas tem uma tremenda vantagem de sobrevivência sobre os que não têm essa capacidade.

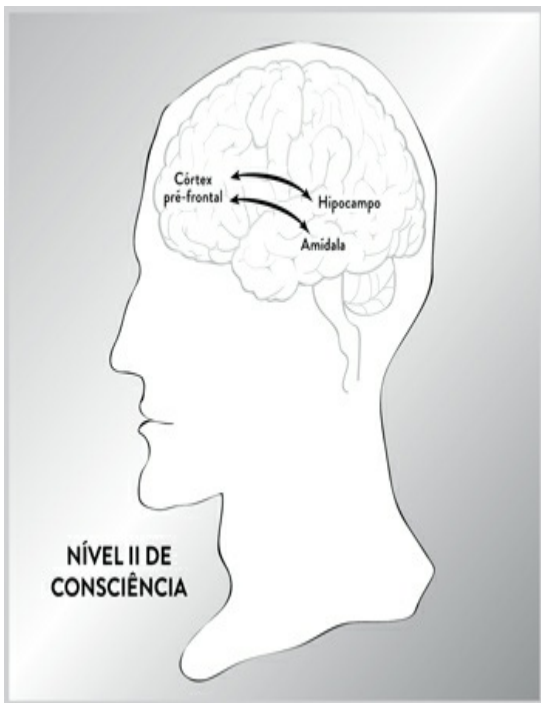


Figura 8. As emoções se originam e são processadas no sistema límbico. No Nível II de consciência, somos continuamente bombardeados com informações sensoriais, e as emoções são respostas rápidas a emergências, dadas pelo sistema límbico, que não precisam de autorização do córtex pré-frontal. O hipocampo também é importante no processamento da memória. Assim, o Nível II de consciência, em seu cerne, envolve a reação da amígdala, do hipocampo e do

córtex pré-frontal.

A Teoria da Mente permite formar alianças com outros, isolar inimigos e consolidar amizades, o que aumenta bastante o poder e as chances de sobrevivência e acasalamento. Alguns antropólogos acreditam até que o domínio da Teoria da Mente foi essencial para a evolução do cérebro.

Mas como se chega à Teoria da Mente? Uma pista foi dada em 1996, com a descoberta dos “neurônios espelho” pelos drs. Giacomo Rizzolatti, Leonardo Fogassi e Vittorio Gallese. Esses neurônios são ativados quando executamos uma tarefa e vemos alguém executando a mesma tarefa. Os neurônios espelho disparam tanto para emoções como para atos físicos. Se sentimos algo e achamos que outra pessoa está tendo a mesma emoção, os neurônios espelho disparam.

Os neurônios espelho são essenciais para a imitação e para a empatia, dando não só a capacidade de copiar tarefas complexas efetuadas por outros, mas também de sentir as emoções que outros devem estar sentindo. Assim, os neurônios espelho provavelmente foram essenciais para nossa evolução como seres humanos, pois a cooperação é essencial para manter a tribo unida.

Os neurônios espelho foram descobertos primeiro nas áreas pré-motoras de cérebros de macacos, e então foram encontrados no córtex pré-frontal dos humanos. O dr. V. S. Ramachandran acredita que os neurônios espelho foram essenciais para nosso poder de autoconsciência, e conclui: “Prevejo que os neurônios espelho serão para a psicologia o que o DNA é para a biologia: vão fornecer uma estrutura unificada e ajudar a explicar uma série de capacidades mentais que até então permaneceram misteriosas e inacessíveis a experimentos.” (Observamos, porém, que todo resultado científico tem que ser testado e reconfirmado. Não há dúvida de que certos neurônios desempenham esse papel crucial ligado a empatia, imitação etc., mas há um debate sobre a identidade desses neurônios espelho. Por exemplo, alguns críticos alegam que talvez esses comportamentos sejam comuns a vários neurônios, e que não existe uma classe específica de neurônios dedicados a esses comportamentos.)

NÍVEL III: SIMULANDO O FUTURO

O nível mais alto de consciência, associado originalmente ao *Homo sapiens*, é o Nível III, onde tomamos nosso modelo do mundo e a partir dele fazemos simulações para o futuro. Analisamos lembranças de pessoas e eventos, depois simulamos o futuro fazendo várias conexões causais e formamos uma árvore “causal”. Ao olharmos os diversos rostos numa festa, nos fazemos perguntas

simples: Como tal pessoa pode me ajudar? Que efeito terão no futuro as fofocas percorrendo o salão? Tem alguém contra mim?

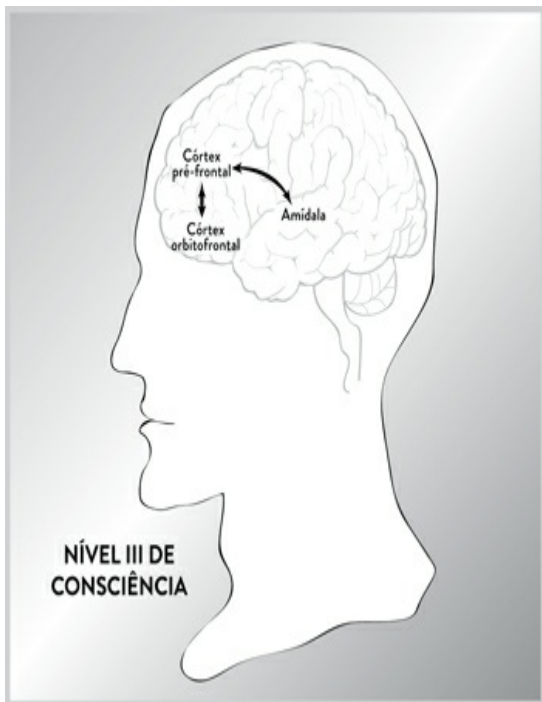


Figura 9. A simulação do futuro, o cerne do Nível III de consciência, é mediada pelo córtex pré-frontal dorsolateral, o diretor do cérebro, com a competição entre o centro de prazer e o córtex orbitofrontal (que age verificando nossos

impulsos). Isso lembra um pouco a ideia de Freud da luta entre a consciência e os desejos. O verdadeiro processo de simulação do futuro ocorre quando o córtex pré-frontal acessa as lembranças do passado a fim de delinear eventos futuros.

Digamos que você acabou de perder o emprego e está procurando outro desesperadamente. Nesse caso, enquanto você conversa com várias pessoas na festa, sua mente está fervilhando, simulando o futuro com cada pessoa com quem conversa. Você se pergunta: Como posso causar boa impressão a essa pessoa? Que assuntos devo abordar para fazer bonito? Será que ela pode me oferecer um emprego?

Varreduras cerebrais recentes ajudaram a esclarecer como o cérebro simula o futuro. Essas simulações são feitas principalmente no córtex pré-frontal dorsolateral, o diretor-geral do cérebro, usando lembranças do passado. Por um lado, as simulações do futuro podem produzir resultados desejáveis e prazerosos, e nesse caso os centros de prazer se acendem (no núcleo accumbens e no hipotálamo). Por outro lado, se esses resultados também têm aspectos negativos, o córtex orbitofrontal se manifesta para advertir sobre possíveis perigos. Então, entre diferentes partes do cérebro, há um conflito concernente ao futuro, que pode ter resultados desejáveis ou indesejáveis. Por fim, é o córtex pré-frontal dorsolateral que faz a mediação entre eles e toma a decisão final. (Ver Figura 9.) Alguns neurologistas observaram que esse conflito lembra, de forma aproximada, a dinâmica freudiana do ego, id e superego.

O MISTÉRIO DA AUTOCONSCIÊNCIA

Se a teoria do espaço-tempo da consciência está correta, também nos dá uma definição precisa de autoconsciência. Em vez de referências vagas e circulares, precisamos de uma definição verificável e útil. Vamos definir a autoconsciência da seguinte maneira:

Autoconsciência é criar um modelo do mundo e simular um futuro no qual estaremos.

Portanto, os animais têm alguma autoconsciência, dado que, para sobreviver e se acasalar, eles precisam saber onde estão situados, mas isso é muito limitado pelo instinto.

A maioria dos animais, quando colocados diante de um espelho, ignora ou ataca, sem descobrir que a imagem é deles mesmos. (Isso é chamado de “teste do espelho”, que remonta a Darwin.) No entanto, animais como os elefantes, os grandes macacos, várias espécies de golfinhos, orcas e pega-rabudas conseguem

entender que a imagem que vemos no espelho é a deles.

Os humanos, porém, dão um grande passo à frente e fazem constantemente simulações do futuro em que aparecem como atores principais. Sempre enfrentamos situações – ir a um encontro, nos candidatar a um emprego, mudar de carreira – que não são determinadas pelo instinto. É extremamente difícil impedir nosso cérebro de simular o futuro, embora alguns métodos elaborados (meditação, por exemplo) tenham sido utilizados para se conseguir isso.

O devaneio, por exemplo, geralmente consiste numa atuação em diferentes futuros possíveis para atingir uma meta. Como nos orgulhamos de conhecer nossos pontos fortes e limitações, temos facilidade de nos colocarmos no modelo e apertar o botão “play” para atuar em cenários hipotéticos, como um ator numa peça virtual.

ONDE “EU” ESTOU?

Provavelmente, há uma parte específica do cérebro cuja tarefa é unificar os sinais vindos dos dois hemisférios para criar uma ideia de nós mesmos. O dr. Todd Heatherton, psicólogo do Dartmouth College, acredita que essa região se localiza no córtex pré-frontal, no chamado córtex pré-frontal medial. O biólogo dr. Carl Zimmer escreve que “o córtex pré-frontal medial pode desempenhar o mesmo papel para o ‘eu’ que o hipocampo desempenha para a memória, (...) pode estar sempre tecendo uma ideia de quem nós somos”. Em outras palavras, pode ser o portal para o conceito de “eu”, a região central do cérebro que funde, integra e inventa uma narrativa unificada de quem nós somos. (Isso não significa, porém, que o córtex pré-frontal medial seja o homúnculo sentado no cérebro controlando tudo.)

Se essa teoria for verdadeira, então o cérebro em sossego, quando estamos ociosamente devaneando sobre nossos amigos e nós mesmos, ficaria mais ativo do que o normal, mesmo quando outras partes das regiões sensoriais cerebrais estão quietas. De fato, a varredura cerebral mostra isso. O dr. Heatherton conclui que “na maior parte do tempo, devaneamos – pensamos sobre alguma coisa que nos aconteceu, ou sobre o que achamos de outras pessoas. Isso envolve autorreflexão”.

A teoria do espaço-tempo diz que a consciência é influenciada por várias subunidades do cérebro, cada uma competindo com as outras para criar um modelo do mundo e, mesmo assim, nossa consciência parece uniforme e contínua. Como isso é possível, já que todos nós achamos que nosso “eu” é ininterrupto e está sempre no comando?

No capítulo anterior, vimos a difícil condição de pacientes com cérebro

dividido lutando com as duas mãos que, literalmente, têm mentes próprias. Parece que existem dois centros de consciência vivendo no mesmo cérebro. Sendo assim, por que criamos a ideia de que temos um “eu” unificado, coeso, existindo em nosso cérebro?

Perguntei a uma pessoa que poderia ter a resposta, o dr. Michael Gazzaniga, que passou várias décadas estudando o estranho comportamento de pacientes com cérebro dividido. Ele notou que o cérebro esquerdo desses pacientes, quando confrontado com o fato de que parecia haver dois centros de consciência residindo no mesmo crânio, procurava dar estranhas explicações, por mais bobas que fossem. Ele me disse que, diante de um paradoxo óbvio, o cérebro esquerdo “fabulava” uma resposta para explicar fatos inconvenientes. O dr. Gazzaniga acredita que isso nos dá uma falsa ideia de que somos unificados e inteiros. Ele chama o cérebro esquerdo de “intérprete”, que fica inventando ideias para tamponar inconveniências e lacunas em nossa consciência.

Por exemplo: num experimento, ele mostrou em flash a palavra “vermelho” só para o cérebro esquerdo de um paciente, e a palavra “banana” só para o cérebro direito. (Notem que o cérebro esquerdo, dominante, não está sabendo da “banana”.) Depois ele pediu ao sujeito que pegasse uma caneta com a mão esquerda (que é governada pelo cérebro direito) e fizesse um desenho. Naturalmente, ele desenhou uma banana. Lembrem que o cérebro direito podia fazer isso porque tinha visto a banana, mas o cérebro esquerdo não sabia que a banana tinha sido mostrada para o cérebro direito.

Ele perguntou ao paciente por que tinha desenhado a banana. Como apenas o lado esquerdo controla a fala, e como o cérebro esquerdo não sabia da banana, o paciente deveria ter respondido “não sei”. Em vez disso, ele falou: “É mais fácil desenhar com essa mão porque desliza melhor.” O dr. Gazzaniga observou que o cérebro esquerdo estava tentando achar uma desculpa para esse fato inconveniente, pois o paciente não sabia por que havia desenhado uma banana.

Ele concluiu que “é o hemisfério esquerdo que gera a tendência de encontrar ordem no caos, que tenta encaixar tudo numa história e pôr num contexto. Parece que ele é direcionado para criar hipóteses sobre a estrutura do mundo, mesmo diante da evidência de que não existe padrão nenhum”.

É daí que vem nossa ideia de um “eu” unificado. Embora a consciência seja uma colcha de retalhos de tendências rivais e frequentemente contraditórias, o cérebro esquerdo ignora as incoerências e tampona lacunas óbvias a fim de nos dar um senso homogêneo de um “eu” único. Em outras palavras, o cérebro esquerdo está sempre arrumando desculpas, algumas descabidas e absurdas, para dar sentido ao mundo. Está sempre perguntando “por quê?” e inventando desculpas, mesmo quando a pergunta não tem resposta.

Há, provavelmente, uma razão evolutiva para termos um cérebro dividido em dois. Um diretor-geral com bastante experiência estimula seus auxiliares a tomar

posições opostas numa questão, para promover um debate amplo e ponderado. Muitas vezes a visão correta emerge da densa interação de ideias incorretas. Da mesma forma, as duas metades do cérebro se complementam, oferecendo análises pessimistas/otimistas ou analíticas/holísticas da mesma ideia. Assim, as duas metades do cérebro conseguem defender seus argumentos. De fato, como veremos, certas formas de doença mental podem surgir quando essa interação entre os dois cérebros não funciona.

Agora que já temos uma teoria da consciência, chegou o momento de utilizá-la para entender como a neurociência irá evoluir no futuro. Atualmente, há um vasto e admirável conjunto de experimentos na neurociência que alteram fundamentalmente todo o panorama científico. Utilizando o poder do eletromagnetismo, os cientistas agora conseguem investigar os pensamentos das pessoas, enviar mensagens telepáticas, controlar telecineticamente os objetos à volta, gravar lembranças, e talvez aprimorar nossa inteligência.

Talvez a aplicação mais imediata e prática dessa nova tecnologia seja algo que já foi considerado totalmente impossível: a telepatia.

1. *Club* tem o duplo sentido de “clube” e de “taco de golfe, bastão”, daí a piada. (N. da T.)

LIVRO II A MENTE ACIMA DA MATÉRIA

O cérebro, queiram ou não, é uma máquina. Os cientistas chegaram a essa conclusão, não por serem desmancha-prazeres mecanicistas, mas porque reuniram evidências de que todos os aspectos da consciência podem ser ligados no cérebro.

– STEVEN PINKER

3 TELEPATIA – UM DOCE PELO Q UE VOCÊ ESTÁ PENSANDO

Alguns historiadores acreditam que Harry Houdini foi o maior mágico que já existiu. Suas inacreditáveis fugas de caixas trancadas e lacradas e seus números perigosos deixavam o público sem fôlego. Ele fazia pessoas desaparecerem e reaparecerem nos lugares mais inesperados. E lia a mente das pessoas.

Ou pelo menos é o que parecia.

Houdini sofria para explicar que tudo o que ele fazia era ilusão, uma série de truques de prestidigitação. Ele ressaltava que ler a mente era impossível. E ficava tão indignado com mágicos inescrupulosos enganando ricos patrocinadores com truques baratos e sessões espíritas, que viajou pelo país expondo as enganações e afirmando que era capaz de reproduzir qualquer fenômeno mental alegado pelos charlatões. Chegou a fazer parte de um comitê organizado pela revista *Scientific American*, que oferecia um generoso prêmio a quem provasse que possuía poderes mentais. (Ninguém jamais recebeu o prêmio.)

Houdini acreditava que a telepatia era impossível. Mas a ciência está provando que Houdini estava errado.

A telepatia é hoje tema de intensas pesquisas em universidades de todo o mundo, onde cientistas já são capazes de usar sensores avançados para ler palavras isoladas, imagens e pensamentos no cérebro das pessoas. Isso pode alterar o modo de comunicação com vítimas de derrame e acidentes que ficam “aprisionadas” no corpo, incapazes de manifestar seus pensamentos exceto piscando os olhos. Mas isso é apenas o começo. A telepatia pode também mudar radicalmente nossa interação com computadores e com o mundo externo.

De fato, num recente “Next 5 in 5 Forecast”, evento da IBM que prevê cinco desenvolvimentos evolucionários nos próximos cinco anos, cientistas da empresa afirmaram que seremos capazes de nos comunicar mentalmente com computadores, talvez substituindo o mouse e os comandos de voz. Isso significa usar o poder da mente para telefonar para alguém, pagar contas com cartão de crédito, dirigir carro, marcar compromissos, criar belas sinfonias e obras de arte etc. As possibilidades são infinitas, e parece que tudo – de gigantes da computação, educadores, empresas de videogame, estúdios de música, até o Pentágono – está convergindo para essa tecnologia.

A verdadeira telepatia, encontrada na ficção científica e em histórias fantásticas, não é possível sem auxílio externo. Como sabemos, o cérebro é elétrico. Em geral, quando um elétron é acelerado, emite radiação eletromagnética. O mesmo se aplica aos elétrons oscilando no cérebro, que emitem ondas de rádio. Mas esses sinais são muito fracos para serem detectados por outras pessoas, e ainda que pudéssemos perceber essas ondas de rádio, seria difícil entendê-las. A evolução não nos deu a capacidade de decifrar esse conjunto de sinais de rádio aleatórios, mas os computadores conseguem. Os

cientistas são capazes de fazer aproximações grosseiras do pensamento, usando a varredura por EEG. O sujeito coloca um capacete com sensores de EEG e se concentra em determinada imagem – digamos, a figura de um carro. Os sinais do EEG são gravados a cada imagem e assim criam um dicionário rudimentar, com correspondência de um a um entre os pensamentos da pessoa e a imagem do EEG. Depois, quando mostram à pessoa a figura de outro carro, o computador reconhece o padrão do EEG correspondente a um carro.

A vantagem dos sensores de EEG é que são rápidos e não são invasivos. Basta colocar um capacete contendo vários eletrodos direcionados à superfície do cérebro e o EEG identifica rapidamente os sinais que mudam a cada milissegundo. Mas o problema com os sensores de EEG, como vimos, é que as ondas eletromagnéticas se deterioram quando atravessam o crânio, e é difícil localizar a origem exata das ondas. Esse método pode dizer se você está pensando num carro ou numa casa, mas não consegue recriar a imagem do carro. É aí que entra o trabalho do dr. Jack Gallant.

VÍDEOS DA MENTE

O epicentro de grande parte dessa pesquisa é a Universidade da Califórnia em Berkeley, onde completei meu doutorado em física teórica, anos atrás. Tive o prazer de conhecer o laboratório do dr. Gallant, cuja equipe tinha realizado uma proeza considerada impossível: gravar em fita de vídeo os pensamentos de alguém. “Isso é um salto muito importante para a reconstrução de imagens internas. Estamos abrindo uma janela para ver os filmes da nossa mente”, disse Gallant.

Quando visitei o laboratório, a primeira coisa que notei foi a equipe de jovens pós-doutores e alunos de pós-graduação, todos muito interessados, amontoados diante das telas dos computadores, olhando atentamente para a reconstrução das imagens de uma varredura cerebral. Conversando com a equipe de Gallant, a sensação é de estar testemunhando a construção da história científica.

Gallant me explicou que o indivíduo se deita numa maca que é lentamente introduzida, começando pela cabeça, num aparelho imenso, o mais avançado de IRM, que custa mais de 3 milhões de dólares. Ele vê então vários clipes de filmes (como os trailers do YouTube). Para acumular dados suficientes, é preciso ficar imóvel durante horas assistindo aos clipes, uma tarefa realmente árdua. Perguntei a um pós-doutor, dr. Shinji Nishimoto, como eles encontravam voluntários dispostos a passar horas e horas imóveis, com apenas fragmentos de vídeos para ocupar o tempo. Ele disse que o pessoal do laboratório, alunos de pós-graduação e pós-doutores, se voluntariavam como cobaias da sua própria

pesquisa.

Enquanto assistem aos cliques, o aparelho de IRM cria uma imagem em 3D do fluxo sanguíneo no cérebro. A imagem produzida pela ressonância magnética parece um grande conjunto de trinta mil pontos, ou vóxeis. Cada vóxel representa um pontinho de energia neural, e a cor do ponto corresponde à intensidade do sinal e ao fluxo do sangue. Pontos vermelhos representam locais de grande atividade neural, e pontos azuis representam locais de atividade menor. (A imagem final parece muito com milhares de luzinhas de Natal acesas, na forma de um cérebro. Pode-se ver imediatamente que o cérebro concentra a maior parte de sua energia mental no córtex visual, situado na parte de trás do cérebro, enquanto assiste aos vídeos.)

O aparelho de IRM de Gallant é tão potente que identifica de duzentas a trezentas regiões do cérebro e mostra imagens instantâneas com cem pontos, em média, em cada região do cérebro. (Um objetivo das futuras gerações de IRM é fornecer uma resolução ainda mais nítida, aumentando o número de pontos por região do cérebro.)

A princípio, esse amontoado de pontos coloridos em 3D parece uma confusão sem nexos. Mas após anos de pesquisa, Gallant e sua equipe desenvolveram uma fórmula matemática que começa a encontrar relações entre certos aspectos da imagem (bordas, texturas, intensidade etc.) e os vóxeis da IRM. Por exemplo: se olharmos para uma linha divisória, notamos que separa áreas mais claras e mais escuras, e a borda gera um certo padrão de vóxeis. E depois de muitos voluntários assistirem a uma grande quantidade de vídeos, essa fórmula matemática é aprimorada, permitindo ao computador analisar como todos esses tipos de imagens são convertidos em vóxeis de IRM. A certa altura, os cientistas conseguiram constatar uma correlação direta entre certos padrões de vóxeis de IRM e aspectos dentro de cada imagem.

Nesse momento, o voluntário assiste a outro trailer. O computador analisa os vóxeis gerados durante a exibição do vídeo e recria uma imagem aproximada da original. Entre os cem cliques, o computador seleciona as imagens mais parecidas com as que o indivíduo acabou de ver e então junta imagens para criar uma aproximação melhor. Desse modo, o computador é capaz de criar um vídeo difuso das imagens que passam pela mente do sujeito. A fórmula matemática do dr. Gallant é tão versátil que pode tomar um grupo de vóxeis de IRM e convertê-los numa imagem, ou fazer o inverso, tomar uma imagem e convertê-la em vóxeis de IRM.

Tive a oportunidade de assistir a um vídeo criado pela equipe do dr. Gallant, e foi muito impressionante. Era como assistir com óculos escuros a um filme com rostos, animais, cenas de rua e prédios. Mesmo sem conseguir ver detalhes de cada rosto ou de cada animal, era possível identificar claramente o tipo de objeto.

Esse programa consegue decodificar não só o que está sendo visto, como também as figuras imaginárias circulando na cabeça. Digamos que peçam a você para pensar na *Mona Lisa*. Sabemos, a partir de varreduras por IRM, que, mesmo que alguém não esteja vendo o quadro com os próprios olhos, o córtex visual do cérebro se acende. O programa do dr. Gallant digitaliza o cérebro enquanto o voluntário está pensando na *Mona Lisa* e depois pesquisa arquivos de dados de imagens, tentando achar a imagem mais aproximada. No experimento a que assisti, o computador selecionou uma imagem da atriz Selma Hayek como a melhor aproximação da *Mona Lisa*. Certamente, qualquer pessoa consegue reconhecer com facilidade centenas de rostos, mas o fato de o computador analisar uma imagem dentro do cérebro de alguém e escolher essa imagem, dentre milhões de outras disponíveis, é impressionante.

A finalidade de todo esse processo é criar um dicionário preciso que permita fazer rapidamente a correspondência entre um objeto do mundo real e o padrão da IRM do cérebro. Em geral, é muito difícil obter uma correspondência detalhada, o que pode levar anos, mas algumas categorias são realmente fáceis de ler, bastando pesquisar fotografias. O dr. Stanislas Dehaene, do Collège de France, em Paris, examinava varreduras por IRM do lobo parietal – onde os números são reconhecidos –, quando um pós-doutor da equipe disse que apenas olhando o padrão de IRM ele podia dizer qual número o sujeito estava vendo. De fato, certos números criavam padrões distintos no mapa de IRM. Ele observa: “Tomando duzentos vóxeis nessa área e vendo quais estão ativos e quais estão inativos, pode-se construir um equipamento com a capacidade de decodificar o número que está retido na memória.”

Ainda não sabemos quando conseguiremos obter vídeos de boa qualidade dos nossos pensamentos. Infelizmente, perde-se informação quando a pessoa está visualizando uma imagem. Isso é corroborado pela varredura cerebral. Quando se compara o mapa de IRM de um cérebro que está olhando para uma flor com um mapa de IRM de um cérebro que está pensando numa flor, vê-se imediatamente que a segunda imagem tem muito menos pontos que a primeira. Assim, embora essa tecnologia vá se desenvolver muito nos próximos anos, jamais será perfeita. (Certa vez li um conto em que um homem encontra um gênio que se oferece para criar qualquer coisa que o homem imaginar. O homem pede um carro de luxo, um jatinho e um milhão de dólares. A princípio, o homem fica em êxtase. Mas quando ele vê de perto, descobre que o carro e o avião não têm motores, e que a imagem do dinheiro está toda borrada. Nada presta. Isso acontece porque nossas lembranças são apenas aproximações das coisas reais.)

Mas, em vista da rapidez com que os cientistas estão começando a decodificar os padrões de IRM no cérebro, será que em breve conseguiremos realmente ler palavras e pensamentos circulando na mente?

LENDO A MENTE

Num prédio ao lado do laboratório de Gallant, o dr. Brian Pasley e seus colegas estão literalmente lendo pensamentos – pelo menos em princípio. Uma pós-doutora da equipe, Sara Szczepanski, me explicou como eles são capazes de ler palavras na mente.

Os cientistas usaram a tecnologia de ECOG (eletrocorticograma), que é um grande avanço em relação à confusão de sinais produzidos pelo EEG. A varredura por ECOG é a mais avançada em termos de precisão e resolução, pois os sinais são gravados diretamente do cérebro, sem passar pelo crânio. O problema é que é preciso remover uma porção do crânio para colocar uma malha, contendo 64 eletrodos dispostos numa grade 8x8, diretamente no cérebro exposto.

Por sorte, eles conseguiram autorização para conduzir experimentos com o ECOG em pacientes epiléticos que sofriam de convulsões debilitantes. A malha do ECOG foi colocada no cérebro dos pacientes durante uma cirurgia cerebral realizada por médicos na Universidade da Califórnia em São Francisco, ali perto.

Enquanto os pacientes ouvem diversas palavras, os sinais emitidos pelo cérebro passam pelos eletrodos e são gravados. Fizeram então um dicionário com a correspondência entre as palavras e os sinais que emanavam dos eletrodos colocados no cérebro. Quando uma palavra é falada, vê-se o mesmo padrão elétrico. Essa correspondência significa também que, se alguém está pensando em determinada palavra, o computador pode captar os sinais característicos e identificar a palavra.

Essa tecnologia possibilita ter uma conversa inteiramente telepática. E vítimas de derrame que ficam totalmente paralisadas podem “falar” através de um sintetizador de voz que reconhece os padrões cerebrais de cada palavra.

Não é de surpreender que a interface cérebro-máquina (ICM) tenha se tornado um ótimo campo de exploração, com laboratórios em todo o país fazendo descobertas significativas. Resultados semelhantes foram obtidos por cientistas da Universidade de Utah em 2011. Eles colocaram grades, com 16 eletrodos cada uma, no córtex motor facial (que controla os movimentos da boca, lábios, língua e face) e na área de Wernicke, que processa informações sobre linguagem.

Pediram ao paciente que dissesse dez palavras comuns, como “sim” e “não”, “quente” e “frio”, “sede” e “fome”, “olá” e “tchau”, “mais” e “menos”. Usando um computador para gravar os sinais cerebrais quando as palavras eram pronunciadas, os cientistas criaram uma correspondência básica um para um entre as palavras e os sinais computacionais emitidos pelo cérebro. Depois, quando o paciente dizia certas palavras, eles conseguiam identificar cada uma com uma precisão de 76% a 90%. O passo seguinte é usar grades com 121

eletrodos para se obter uma resolução melhor.

No futuro, esse procedimento poderá ser útil para indivíduos com derrame ou doenças paralisantes, como a doença de Lou Gehrig, que poderão falar usando a técnica do cérebro ao computador.

DIGITANDO COM A MENTE

Na Mayo Clinic de Minnesota, o dr. Jerry Shih conectou pacientes epiléticos a sensores de ECOG para que aprendessem a digitar com a mente. A calibragem desse equipamento é simples. Mostra-se ao paciente uma série de letras para que ele se concentre mentalmente em cada uma delas. Um computador registra os sinais emitidos pelo cérebro a cada letra observada. Como ocorreu com outros experimentos, uma vez criado o dicionário um para um, a tarefa é simples e basta a pessoa pensar numa letra, que essa letra aparece numa tela, por meio apenas do poder da mente.

O dr. Shih, líder desse projeto, diz que a precisão dessa máquina é de quase 100%. Ele acredita que poderá criar um aparelho para gravar imagens, e não somente palavras, que os pacientes tenham na mente. Isso pode ter aplicação para artistas e arquitetos, mas o grande problema da tecnologia de ECOG, como vimos, é ter que expor o cérebro da pessoa.

Enquanto isso, os digitadores por EEG, por não serem invasivos, estão entrando no mercado. Não têm o rigor e a precisão dos digitadores por ECOG, mas possuem a vantagem de poder ser vendidos em lojas. A Guger Technologies, sediada na Austrália, demonstrou recentemente um digitador por EEG numa feira comercial. Segundo seus representantes, as pessoas podem aprender a operar a máquina em cerca de dez minutos, e podem digitar uma média de cinco a dez palavras por minuto.

TEXTO E MÚSICA POR TELEPATIA

O próximo passo pode ser a transmissão de conversas inteiras, o que pode acelerar muito a transmissão telepática. O problema, porém, é que exige montar um mapa de correspondências um para um entre milhares de palavras e os sinais de EEG, IRM ou ECOG. Mas se alguém puder, por exemplo, identificar os sinais de várias centenas de palavras selecionadas, será capaz de transmitir rapidamente as palavras mais encontradas numa conversa comum. Isso significa que poderemos pensar em palavras para formar frases ou parágrafos inteiros de

uma conversa, e o computador as digitará.

Seria extremamente útil para jornalistas, ensaístas, romancistas e poetas, que só precisariam pensar e o computador captaria o ditado pensado. O computador seria também uma espécie de secretário mental. Seria possível dar a esse robô secretário instruções sobre um jantar, uma viagem de avião, ou férias, e ele se encarregaria de tomar as providências e fazer as reservas.

Não só textos, mas também a música poderá um dia ser transcrita por esse meio. Os músicos só precisariam entoar mentalmente uma melodia e o computador faria a transcrição, em notação musical. Para isso, a pessoa entoa uma série de notas, o que gera certos sinais elétricos para cada uma. Mais uma vez, cria-se um dicionário, de modo que, quando a pessoa pensa numa nota musical, o computador a transcreve em notação musical.

Na ficção científica, os telepatas geralmente se comunicam superando as barreiras da língua, porque os pensamentos são considerados universais. No entanto, não é bem assim. Sentimentos e emoções podem ser não verbais e universais, de modo que se pode enviá-los telepaticamente a qualquer um, mas o pensamento racional está tão ligado à linguagem que é muito improvável que pensamentos complexos ultrapassem a barreira da língua. As palavras ainda seriam enviadas telepaticamente no idioma original.

CAPACETES TELEPÁTICOS

Na ficção científica, encontramos também muitos capacetes telepáticos. É só colocar e – pronto! – conseguimos ler a mente de outras pessoas. De fato, o Exército dos Estados Unidos manifestou interesse por essa tecnologia. Sob fogo cerrado, com bombas explodindo por todo lado e balas zunindo sobre a cabeça, um capacete telepático poderia salvar vidas, pois é difícil transmitir ordens no meio do barulho e da fúria do campo de batalha. (Posso dar testemunho disso. Anos atrás, durante a Guerra do Vietnã, servi na Infantaria dos Estados Unidos no Forte Benning, nos arredores de Atlanta, na Geórgia. Em exercícios com armas, o som de rajadas de balas e granadas explodindo perto de mim era ensurdecedor. Era tão intenso que eu não ouvia mais nada. Depois fiquei com um zumbido alto nos ouvidos que durou três dias inteiros.) Tendo um capacete telepático, o soldado poderia se comunicar mentalmente com seu pelotão em meio aos estrondos.

Recentemente, o Exército concedeu 6,3 milhões de dólares ao dr. Gerwin Schalk, do Albany Medical College, mesmo sabendo que um capacete telepático totalmente funcional está a anos de distância. Schalk faz experimentos com a tecnologia ECOG que, como vimos, exige que se coloque uma malha de eletrodos diretamente no cérebro exposto. Com esse método, seus computadores

conseguiram reconhecer as vogais e 36 palavras num cérebro pensante. Em alguns experimentos, ele chegou perto de 100% de precisão. Mas atualmente isso ainda é impraticável para o Exército dos Estados Unidos, pois exige a remoção de parte do crânio no ambiente limpo e esterilizado de um hospital. Além disso, reconhecer vogais e meia dúzia de palavras está muito longe de enviar mensagens urgentes para o quartel-general sob fogo cruzado. Mas seus experimentos com o ECOG demonstraram que é possível se comunicar mentalmente num campo de batalha.

Outro método está sendo explorado pelo dr. David Poeppel, da Universidade de Nova York. Em vez de abrir o crânio, ele emprega a tecnologia MEG, usando pequenas explosões de energia magnética em vez de eletrodos para criar cargas elétricas no cérebro. Além de não invasiva, a vantagem da tecnologia MEG é medir com precisão a fugaz atividade neural, em contraste com a varredura por IRM, que é mais lenta. Em seus experimentos, Poeppel conseguiu registrar atividade elétrica no córtex auditivo quando a pessoa pensa numa certa palavra. A desvantagem é que essa gravação ainda exige máquinas grandes, do tamanho de uma mesa, para gerar um pulso magnético.

Obviamente, queremos um método não invasivo, portátil e preciso. O dr. Poeppel espera que seu trabalho com a tecnologia MEG complemente o que vem sendo feito com sensores de EEG. Mas a verdadeira telepatia com capacetes ainda está a anos de distância, porque as varreduras, tanto por MEG como por EEG, têm pouca precisão.

IRM NO TELEFONE CELULAR

Hoje, ainda somos refreados pela natureza relativamente tosca dos instrumentos existentes. Mas, com o passar do tempo, equipamentos cada vez mais sofisticados irão investigar a mente mais a fundo. A próxima grande virada deve ser IRM portátil.

O motivo pelo qual os equipamentos atuais de IRM são tão grandes é a necessidade de um campo magnético uniforme para se obter uma boa resolução. Quanto maior o ímã, mais uniforme é a captação, e maior é a precisão das imagens finais. Entretanto, os físicos conhecem as propriedades matemáticas exatas de campos magnéticos (foram definidas pelo físico James Clerk Maxwell nos anos 1860). Em 1993, na Alemanha, o dr. Bernhard Blümich e seus colegas criaram o menor aparelho de IRM no mundo, do tamanho de uma maleta. Utiliza um campo magnético fraco e distorcido, mas supercomputadores conseguem analisar o campo magnético e fazer as correções, de modo que o aparelho produza imagens realistas em 3D. Como a potência dos computadores é

duplicada a mais ou menos cada dois anos, eles hoje têm potência suficiente para analisar o campo magnético criado pelo aparelho do tamanho de uma maleta e compensar a distorção.

Para demonstrar sua máquina, em 2006 Blümich e seus colegas fizeram varreduras por IRM em Ötzi, o “Homem do gelo”, congelado por cerca de cinco mil e trezentos anos, desde o final da última era do gelo. Como Ötzi foi congelado numa posição peculiar, com os braços afastados do corpo, era difícil enfiá-lo no pequeno cilindro de um aparelho de IRM convencional, mas foi fácil obter fotos com o aparelho de IRM portátil do dr. Blümich.

Esse grupo de físicos estima que, com a potência cada vez maior dos computadores, a máquina de IRM do futuro possa ser do tamanho de um aparelho celular. Os dados básicos gravados no celular serão enviados sem fio para um supercomputador que processará os dados do campo magnético fraco e criará uma imagem em 3D. (A fraqueza do campo magnético será compensada pela maior potência computacional.) Isso irá acelerar muito as pesquisas. “Talvez algo como o tricorder de *Jornada nas estrelas* não esteja mais tão longe”, diz o dr. Blümich. (O tricorder é um pequeno aparelho de mão para varreduras que dá o diagnóstico instantâneo de qualquer doença.) No futuro, possivelmente haverá computadores mais potentes num consultório médico do que hoje num hospital universitário bem equipado. Em vez de esperar autorização de um hospital ou universidade para usar uma máquina de IRM caríssima, será possível obter dados na sala de casa, apenas fazendo uma varredura no corpo com o IRM portátil, e passar os resultados por e-mail para o laboratório analisar.

Isso significa também que, no futuro, talvez consigam fazer um capacete telepático de IRM, com resolução imensamente melhor que um equipamento de EEG. Isso poderá funcionar nas próximas décadas da seguinte maneira: dentro do capacete haverá bobinas eletromagnéticas para produzir um campo magnético fraco e pulsos de rádio sondando o cérebro. Os sinais brutos de IRM são enviados para um computador de bolso preso no cinto de um comandante. As informações são enviadas por rádio para um servidor localizado longe do campo de batalha. O processamento final dos dados é feito por um supercomputador numa cidade distante. A mensagem retorna, por rádio, para os soldados dele no campo de batalha. A tropa ouve a mensagem, por alto-falantes ou por eletrodos colocados no córtex auditivo dos soldados.

DARPA E O APRIMORAMENTO HUMANO

Em vista dos custos de todas essas pesquisas, pode-se perguntar: Quem está pagando? As empresas privadas apenas recentemente demonstraram interesse

nessas tecnologias de ponta, mas ainda é arriscado para muitas delas financiar pesquisas que talvez não deem retorno. Uma das maiores financiadoras é a Darpa, sigla em inglês para Agência de Projetos de Pesquisa Avançada de Defesa, do Pentágono, que esteve à frente de algumas das mais importantes tecnologias do século XX.

A Darpa foi implantada pelo presidente Dwight Eisenhower depois que os russos lançaram o Sputnik em órbita, em 1957, deixando o mundo abismado. Vendo que os Estados Unidos poderiam ser rapidamente ultrapassados em alta tecnologia pela União Soviética, Eisenhower criou a agência, para se manter equiparado aos russos. Com o passar dos anos, os diversos projetos que iniciou cresceram tanto que se tornaram entidades independentes. Uma das primeiras a surgir foi a Nasa.

A estratégia da Darpa parece algo da ficção científica. Sua “única regra é a inovação radical”. A única justificativa para sua existência é “trazer o futuro para o presente”. Os cientistas da agência estão sempre ampliando as fronteiras do que é fisicamente possível. Michael Goldblatt, ex-agente da Darpa, diz que eles tentam não violar as leis da física, “pelo menos, não propositalmente. Ou pelo menos não mais que uma por programa”.

Mas o que separa a Darpa da ficção científica é um histórico realmente espantoso. Um dos primeiros projetos, nos anos 1960, foi a Arpanet, uma rede de telecomunicações de guerra para conectar eletronicamente os cientistas aos agentes do governo durante e depois da Terceira Guerra Mundial. Em 1989, a National Science Foundation decidiu que, em vista do colapso do bloco soviético, não era mais necessário manter segredo sobre aquela tecnologia militar confidencial, e praticamente liberou gratuitamente suas plantas e códigos. A Arpanet acabou se tornando a Internet.

Quando a Força Aérea dos Estados Unidos precisou de um meio de guiar seus mísseis balísticos no espaço, a Darpa ajudou a criar o Projeto 57, altamente secreto, concebido para jogar bombas-H em silos fixos de mísseis soviéticos numa guerra termonuclear. Mais tarde, tornou-se a base do GPS. Em vez de guiar mísseis, hoje guia motoristas desorientados.

A Darpa tem sido protagonista numa série de invenções que alteraram os séculos XX e XXI, responsável inclusive pelos telefones celulares, óculos de visão noturna, avanços nas telecomunicações e satélites meteorológicos. Tive oportunidade de interagir com cientistas e funcionários da Darpa em várias ocasiões. Uma vez almocei com um ex-diretor da agência numa recepção cheia de cientistas e futuristas. Fiz-lhe uma pergunta que sempre me atormentara: Por que temos que confiar em cães para cheirar as bagagens e detectar a presença de explosivos? Certamente, temos sensores capazes de detectar sinais indicativos de substâncias químicas explosivas. Ele respondeu que a Darpa tinha se dedicado ativamente a essa questão, mas sempre havia esbarrado em sérios problemas

técnicos. Ele disse que o olfato dos cães tinha evoluído durante milhões de anos para ser capaz de detectar até um punhado de moléculas, e é extremamente difícil conseguir essa mesma sensibilidade, mesmo com os sensores mais modernos. Num futuro próximo, é provável que continuemos a contar com cães nos aeroportos.

Em outra ocasião, um grupo de físicos e engenheiros da Darpa compareceu a uma palestra minha sobre o futuro da tecnologia. Depois perguntei a eles se tinham preocupações específicas. Disseram que uma questão era a imagem da agência junto ao público. A maioria das pessoas nunca ouviu falar da Darpa, mas algumas a associam a conspirações nefastas do governo, passando pelo acobertamento de OVNI's, Área 51, Caso Roswell, controle do clima etc. Eles lamentaram; se pelo menos os boatos fossem verdadeiros, eles certamente poderiam usar alguma tecnologia alienígena para adiantar muito as pesquisas!

Atualmente, com um orçamento de três bilhões de dólares, a Darpa concentra seus esforços na interface cérebro-máquina. Ao discutir suas possíveis aplicações, o ex-agente da Darpa Michael Goldblatt amplia as fronteiras da imaginação. Ele diz: "Imagine se soldados pudessem se comunicar por meio do pensamento. (...) Imagine a ameaça de um ataque biológico sendo irrelevante. E por um momento considere um mundo em que aprender seja tão fácil quanto comer, e em que a substituição de partes danificadas do corpo seja tão acessível quanto ir a uma lanchonete. Por mais impossíveis que pareçam essas ideias, por mais difíceis que sejam essas tarefas, esse panorama é o trabalho cotidiano do Defense Sciences Office (um ramo da Darpa)."

Na opinião de Goldblatt, os historiadores concluirão que o legado a longo prazo da Darpa será o aprimoramento humano, "nossa força histórica futura". Ele observa que o slogan do Exército americano "Seja Tudo O Que Você Pode Ser" ganha um novo significado quando se contemplam as implicações do aprimoramento humano. Não é à toa que Michael Goldblatt, na Darpa, está se empenhando tanto no aprimoramento humano. Sua filha sofre de paralisia cerebral e passa a vida confinada numa cadeira de rodas. Como exige cuidados externos, a doença lhe impôs limitações, mas ela sempre se colocou acima da adversidade. Frequenta a universidade e sonha em abrir a própria empresa. Goldblatt reconhece que a filha é sua inspiração. Como observou Joel Garreau, editor do *Washington Post*, "o que ele está fazendo é gastar incalculáveis milhões de dólares para criar o que pode ser o próximo passo da evolução humana. E ele sabe que a tecnologia que está ajudando a criar pode algum dia não só permitir à sua filha andar, mas ir além disso".

Ao ouvir falar pela primeira vez em máquinas de leitura da mente, um leigo pode se preocupar com a privacidade. A ideia de uma máquina escondida em algum lugar, lendo seus pensamentos mais íntimos sem permissão, é aflitiva. A consciência humana, como enfatizamos, envolve contínuas simulações do futuro. Para que essas simulações sejam precisas, às vezes imaginamos situações que adentram territórios imorais ou ilegais, e se vamos ou não concretizá-las, preferimos manter em segredo.

Para os cientistas, a vida seria mais fácil se eles pudessem ler a mente das pessoas a distância usando aparelhos portáteis (melhor do que usar um capacete esquisito ou ter o cérebro cirurgicamente exposto), mas as leis da física tornam isso difícil demais.

Quando perguntei ao dr. Nishimoto, que trabalha no laboratório do dr. Gallant em Berkeley, sobre a questão da privacidade, ele sorriu e respondeu que os sinais de rádio se degradam rapidamente fora do cérebro, portanto se tornariam muito difusos e fracos para serem entendidos por qualquer pessoa a mais de poucos metros de distância. Na escola, aprendemos as leis de Newton e que a gravidade diminui proporcionalmente ao quadrado da distância, de modo que duplicando sua distância de uma estrela, o campo gravitacional diminui para um quarto. Mas os campos magnéticos diminuem muito mais do que o quadrado da distância. A maioria dos sinais decresce proporcionalmente ao cubo ou à quarta potência da distância, portanto, se duplicarmos a distância entre o objeto e um aparelho de IRM, o campo magnético cai para a fração de um oitavo, ou menos.

Além disso, haveria interferência do mundo externo mascarando os fracos sinais emitidos pelo cérebro. Esta é uma razão pela qual os cientistas exigem condições controladas de laboratório para trabalhar, e mesmo assim até agora só conseguiram extrair algumas letras, palavras ou imagens de um cérebro pensante. A tecnologia não é adequada para registrar a avalanche de pensamentos que normalmente circulam pela mente enquanto consideramos várias letras, palavras, frases e informações sensoriais. Portanto, usar dispositivos para ler a mente, como se vê no cinema, não é possível hoje e nem será nas próximas décadas.

Num futuro próximo, as varreduras cerebrais continuarão a depender do acesso direto ao cérebro humano em condições de laboratório. Mas, no caso altamente improvável de alguém descobrir uma maneira de ler pensamentos a distância, haverá formas de se defender. Para manter em segredo os pensamentos mais importantes, você poderá usar um escudo para bloquear as ondas cerebrais, evitando que cheguem aonde não devem. Isso pode ser feito com algo chamado Gaiola de Faraday, inventada pelo grande físico inglês Michael Faraday em 1836, embora o efeito tenha sido observado pela primeira vez por Benjamin Franklin. Em linhas gerais, a eletricidade se dispersa rapidamente ao redor de uma gaiola de metal, de modo que o campo elétrico

dentro da gaiola é zero. Para demonstrar isso, muitos físicos (como eu) já entraram numa jaula metálica na qual foram lançadas enormes faíscas elétricas. Milagrosamente, saímos ilesos. É por isso que os aviões não são danificados quando atingidos por raios, e é por isso que cabos condutores são blindados por malhas metálicas. Da mesma forma, um escudo telepático consistirá em uma fina camada de metal colocada em torno da cabeça.

TELEPATIA VIA NANOSSONDAS NO CÉREBRO

Há outro modo de resolver parcialmente a questão da privacidade, bem como a dificuldade de se colocar sensores de ECOG no cérebro. No futuro, talvez seja possível explorar a nanotecnologia, a capacidade de manipular cada átomo, e inserir no cérebro uma rede de nanossondas para captar os pensamentos. As nanossondas podem ser feitas de nanotubos de carbono, que conduzem eletricidade e são tão finos quanto permitem as leis da física atômica. Esses nanotubos são feitos apenas de átomos de carbono arranjados em um tubo com diâmetro de algumas moléculas. Nanotubos são alvo de grande interesse científico e espera-se que, nas próximas décadas, revolucionem as formas usadas pelos cientistas para investigação do cérebro.

As nanossondas são colocadas nas áreas exatas do cérebro encarregadas de determinadas atividades. Para transmitir a fala e a linguagem, elas são colocadas no lobo temporal esquerdo. Para processar imagens visuais, são colocadas no tálamo e no córtex visual. As emoções são enviadas via nanossondas na amígdala e no sistema límbico. Os sinais das nanossondas são enviados para um computador pequeno, que processa os sinais e envia, sem fio, as informações para um servidor, e daí para a internet.

A questão da privacidade estará parcialmente resolvida, pois poderemos controlar quando nossos pensamentos deverão ser enviados por cabos ou pela internet. Qualquer curioso que tenha um receptor pode detectar sinais de rádio, mas não os sinais elétricos enviados por um cabo. O problema de abrir o crânio e o estorvo das ECOG também será resolvido com a inserção de nanossondas por microcirurgia.

Alguns autores de ficção científica conjecturaram que, no futuro, todos os bebês terão implantes indolores de nanossondas ao nascer, de modo que a telepatia fará parte do modo de vida deles. Em *Jornada nas estrelas*, por exemplo, os implantes são sempre colocados nas crianças Borg logo que nascem, para se comunicarem telepaticamente com os outros. Essas crianças nem podem imaginar um mundo em que não exista a telepatia. Presumem que telepatia é a regra.

Como essas nanossondas são minúsculas, seriam invisíveis, não haveria

rejeição social. Embora a sociedade possa repelir a ideia de ter sondas permanentes no cérebro, os autores de ficção científica acham que as pessoas se acostumarão devido à grande utilidade delas, da mesma forma que os bebês de proveta são hoje bem aceitos, apesar da controvérsia inicial a respeito.

QUESTÕES LEGAIS

Num futuro não muito distante, a questão não é se alguém será capaz de ler nossos pensamentos secretamente com um dispositivo remoto escondido, mas se vamos permitir que nossos pensamentos sejam registrados. Então, o que acontecerá se alguém sem escrúpulos tiver acesso, sem autorização, aos nossos arquivos? Isso levanta uma questão ética, pois não queremos que nossos pensamentos sejam lidos contra nossa vontade. O dr. Brian Pasley diz: “Há questões éticas, não com relação à pesquisa em curso, mas com seus possíveis desdobramentos. É preciso haver um equilíbrio. Se, de algum modo, pudermos decodificar instantaneamente os pensamentos de alguém, isso poderá trazer grandes benefícios para milhares de pessoas com danos graves e que não estão conseguindo se comunicar nesse momento. Por outro lado, há uma grande preocupação de que isso possa ser aplicado a quem não quer.”

Quando for possível ler a mente e fazer gravações, vão surgir inúmeras questões éticas e legais. Isso acontece sempre que é introduzida uma nova tecnologia. Historicamente, costuma levar anos até que se faça uma lei para dar conta de todas as implicações.

Por exemplo: as leis de direitos autorais terão que ser refeitas. O que acontece se alguém roubar sua invenção porque leu seus pensamentos? Vai ser possível patentear pensamentos? A quem realmente pertence a ideia?

Outro problema ocorre se o governo estiver envolvido. Como disse John Perry Barlow, poeta e letrista da banda Grateful Dead: “Confiar no governo para proteger sua privacidade é o mesmo que pedir a um *voyeur* para instalar persianas nas janelas da sua casa.” A polícia terá autoridade para ler seus pensamentos num interrogatório? Já existem casos na justiça com pareceres sobre suspeitos que se recusam a se submeter a um exame de DNA. No futuro, o governo poderá ler seus pensamentos sem a sua autorização? E se puder, eles serão lícitos num tribunal de justiça? Até que ponto serão confiáveis? Da mesma forma que os detectores de mentiras por IRM medem apenas o aumento da atividade cerebral, é importante notar que pensar sobre um crime e realmente cometer um crime são duas coisas muito diferentes. Durante o julgamento, o advogado de defesa pode argumentar que eram apenas pensamentos ao acaso, fantasias, e nada mais.

Outra área cinzenta concerne aos direitos de pessoas com paralisia. Se quiserem fazer um testamento, uma declaração, a varredura cerebral será suficiente para lavrar um documento legal? Suponhamos que uma pessoa totalmente paralisada tenha a mente ágil, em plena atividade, e queira assinar um contrato ou gerenciar suas finanças. Esses documentos terão valor legal, sabendo-se que a tecnologia pode não ser perfeita?

Não há lei da física para resolver tais questões éticas. Em última instância, quando as tecnologias tiverem amadurecido, essas questões terão que ser decididas no tribunal, pelo juiz e o júri.

Até lá, governos e corporações terão que inventar meios de impedir espionagem mental. A espionagem industrial já é uma indústria multimilionária, com governos e corporações construindo “salas cofre” caríssimas, cuidadosamente examinadas contra aparelhos de escuta e gravação. No futuro (supondo que seja criado um método de captar ondas cerebrais a distância), as salas cofre terão que ser projetadas de modo que sinais cerebrais não possam vazarem acidentalmente para o mundo externo. Terão paredes metálicas, formando uma Gaiola de Faraday para isolar seu interior.

Cada vez que uma forma de radiação foi explorada, houve tentativas de usá-la para espionagem. As ondas cerebrais provavelmente não serão exceção. O caso mais famoso envolveu um minúsculo aparelho de micro-ondas escondido no Grande Selo dos Estados Unidos na embaixada americana em Moscou. De 1945 a 1952, o aparelho transmitiu mensagens altamente secretas de diplomatas americanos diretamente para os soviéticos. Mesmo durante a Crise de Berlim em 1948 e a Guerra da Coreia, os soviéticos usaram esse grampo para decifrar os planos dos Estados Unidos. Os segredos estariam vazando até hoje, mudando o curso da Guerra Fria e a história mundial, se a artimanha não tivesse sido descoberta acidentalmente por um engenheiro britânico, que ouviu conversas secretas numa estação de rádio aberta. Os engenheiros dos Estados Unidos ficaram horrorizados quando desmontaram o aparelho. Eles não tinham conseguido detectá-lo durante anos porque era passivo, não precisava de nenhuma fonte de energia. (Os soviéticos, muito espertos, conseguiram mantê-lo em segredo porque o aparelho era energizado por feixes de micro-ondas de uma fonte distante.) No futuro, é possível que sejam construídos dispositivos de espionagem para interceptar ondas cerebrais também.

Embora muito dessa tecnologia ainda esteja numa fase primitiva, a telepatia vem se tornando lentamente um fato da vida. No futuro, poderemos interagir com o mundo por meio da mente. Mas os cientistas querem ir além da leitura da mente, porque a leitura ainda pertence à esfera da passividade. Querem ter o papel ativo de mover objetos com a mente. A telecinesia é um poder geralmente atribuído aos deuses. É o poder divino de moldar a realidade segundo sua vontade. É a mais alta expressão de nossos pensamentos e desejos.

Em breve, conseguiremos isso também.

O perigo faz parte do desenvolvimento do futuro. (...) Os principais avanços da civilização são processos que por pouco não arrasaram as sociedades em que ocorrem.

– ALFRED NORTH WHITEHEAD

4 TELECINESIA – A MENTE CONTROLANDO A MATÉRIA

Cathy Hutchinson está presa dentro do próprio corpo. Ela tem paralisia desde os 14 anos, em consequência de um grave derrame. Tetraplégica, ela vive como milhares de pacientes “presos”, que perderam o controle da maioria dos músculos e funções corporais. Passa a maior parte do dia impotente, precisa da assistência de enfermeiros, mas tem a mente lúcida. É uma prisioneira do próprio corpo.

Mas, em maio de 2012, sua sorte mudou radicalmente. Cientistas da Universidade de Brown colocaram na parte superior de seu cérebro um pequeno chip, chamado Braingate, conectado por cabos a um computador. Os sinais do cérebro são transmitidos pela máquina a um braço mecânico robótico. Usando apenas o pensamento, ela aprende gradualmente a controlar os movimentos do braço para conseguir, por exemplo, pegar uma garrafa de água e levá-la à boca. Pela primeira vez, ela é capaz de ter algum controle sobre o mundo ao seu redor.

Como não consegue falar, teve que comunicar seu entusiasmo fazendo movimentos com os olhos. Um dispositivo acompanha seus olhos e traduz os movimentos em mensagens digitadas. Quando lhe perguntaram como se sentia, após anos de aprisionamento numa concha chamada corpo, ela respondeu “extasiada!”. Ansiando pelo dia em que seus outros membros sejam conectados ao cérebro por meio do computador, ela acrescentou: “Eu adoraria ter uma perna robótica.” Antes do derrame, ela gostava muito de cozinhar e cuidar do jardim. “Sei que algum dia isso voltará a acontecer”, ela disse. Devido ao ritmo dos avanços no campo de próteses cibernéticas, seu desejo pode se realizar em breve.

O professor John Donoghue e seus colegas na Universidade de Brown e na Universidade de Utah criaram um minúsculo sensor que funciona como uma ponte para quem não pode mais se comunicar com o mundo externo. Quando o entrevistei, ele me disse: “Pegamos um pequeno sensor, do tamanho de uma aspirina infantil, quatro milímetros, e implantamos na superfície do cérebro. Tem noventa e seis ‘cabelinhos’, ou eletrodos, que captam os impulsos cerebrais. Ele pode captar sinais da intenção da pessoa para mover o braço. Escolhemos o braço por causa de sua importância.” Como o córtex motor vem sendo mapeado minuciosamente há décadas, é possível colocar o chip diretamente sobre os neurônios que controlam membros específicos.

O segredo do Braingate está na conversão dos sinais neurais do chip em comandos definidos para mover objetos, a começar pelo cursor na tela do computador. Donoghue me contou que faz isso pedindo ao paciente para imaginar o movimento do cursor, por exemplo, para a direita. Em poucos minutos, os sinais cerebrais correspondentes ao ato são gravados. Dessa maneira, sempre que o computador detectar e reconhecer um sinal cerebral como aquele,

deve mover o cursor para a direita.

Então, cada vez que essa pessoa pensa em mover o cursor para a direita, o computador obedece. Assim se cria um mapa de correspondências exatas entre certas ações que o paciente imagina e cada ato propriamente dito. O paciente passa a controlar o movimento do cursor geralmente na primeira tentativa.

O Braingate abre a porta para um novo mundo de neuropróteses, permitindo que uma pessoa paralisada movimente membros artificiais com a mente. Além disso, o paciente pode se comunicar diretamente com amigos e parentes. A primeira versão desse chip, testada em 2004, foi concebida para pacientes paraplégicos se comunicarem com um computador portátil. Pouco depois esses pacientes estavam navegando na internet, lendo e escrevendo e-mails, e controlando a cadeira de rodas.

Mais recentemente, uma neuroprótese foi ligada aos óculos do cosmólogo Stephen Hawking. Como o sensor de EEG, esse dispositivo conecta pensamentos a um computador, de modo que ele possa manter contato com o mundo externo. Ainda é bem primitivo, mas com o tempo, os dispositivos com essa finalidade ficarão muito mais sofisticados, sensíveis e com mais canais.

Donoghue me disse que esses avanços podem ter um profundo impacto na vida dos pacientes: “Outra utilidade é que se pode conectar o computador a qualquer aparelho – torradeira, cafeteira elétrica, ar-condicionado, interruptor de luz, máquina de escrever. Hoje é realmente muito fácil fazer isso, e muito barato. Um tetraplégico que não consegue se mexer conseguirá mudar o canal da televisão, acender a luz, e tudo isso sem precisar da ajuda de ninguém.” Virá o tempo em que eles serão capazes de fazer tudo o que faz uma pessoa normal, via computadores.

SOLUÇÃO PARA LESÕES NA MEDULA ESPINHAL

Vários laboratórios estão entrando nessa área. Outra descoberta importante foi feita por cientistas da Northwestern University, que conectaram o cérebro de um macaco ao próprio braço, fazendo um desvio, sem passar pela medula espinhal lesionada. Em 1995, houve o triste caso de Christopher Reeve, que voava pelos ares nos filmes de *Super-Homem* e ficou tetraplégico devido a uma lesão na medula espinhal. Ele teve a infelicidade de cair do cavalo e aterrissar sobre o pescoço, lesionando a coluna logo abaixo da cabeça. Se tivesse vivido mais tempo, poderia ter conhecido o trabalho de cientistas que buscam uma forma de usar o computador para substituir a medula espinhal. Somente nos Estados Unidos, mais de 200 mil pessoas têm algum tipo de lesão na coluna. No passado, essas pessoas teriam morrido pouco depois do acidente, mas, com os avanços no

tratamento de traumas agudos, o número de sobreviventes a essas lesões realmente cresceu nos últimos anos. Também nos afligem as imagens de milhares de soldados vítimas de minas terrestres no Iraque e Afeganistão. E se incluirmos o número de pacientes com paralisias por derrame e outras doenças, como a esclerose lateral amiotrófica (ELA), esse número sobe para dois milhões.

Os cientistas da Northwestern University colocaram um chip de 100 eletrodos diretamente no cérebro de um macaco. Os sinais cerebrais foram minuciosamente registrados quando o macaco pegou uma bola, levantou-a, e a colocou dentro de um tubo. Como cada uma dessas ações corresponde a uma ativação específica de neurônios, os cientistas conseguiram, gradualmente, decodificar os sinais.

Quando o macaco queria mover o braço, os sinais eram processados por um computador usando esse código, e em vez de enviar as mensagens para um braço mecânico, enviavam os sinais diretamente para os nervos do braço do macaco. “Estamos interceptando os sinais elétricos naturais do cérebro, que dizem ao braço e à mão como devem se mover, e enviando esses sinais diretamente para os músculos”, diz o dr. Lee Miller.

Por tentativa e erro, o macaco aprendeu a coordenar os músculos do braço. “Há um processo de aprendizado motor muito semelhante ao que você passa quando usa um computador novo, um mouse diferente, ou uma outra raquete de tênis”, acrescenta o dr. Miller.

É notável que o macaco tenha conseguido dominar tantos movimentos do braço, dado que o chip no cérebro só tinha cem eletrodos. O dr. Miller observa que milhões de neurônios estão envolvidos no controle do braço. O motivo de 100 eletrodos conseguirem gerar uma aproximação razoável do trabalho de milhões de neurônios é que o chip se conecta aos neurônios finais, depois que todo o processamento complexo já foi realizado pelo cérebro. Sem necessidade de uma análise mais sofisticada, os 100 eletrodos são responsáveis simplesmente por fornecer a informação ao braço.

Esse dispositivo é um dos vários que estão sendo desenvolvidos pela universidade, para permitir aos pacientes contornar a lesão na medula espinhal. Outra prótese neural usa o movimento dos ombros para controlar o braço. Encolher os ombros para cima faz a mão se fechar. Abaixar os ombros faz a mão se abrir. O paciente pode também dobrar os dedos para segurar um objeto, como um copo, ou manipular uma chave presa entre o polegar e o indicador.

O dr. Miller conclui: “Essa conexão do cérebro aos músculos poderá um dia ser usada para ajudar pacientes imobilizados por lesões na medula a desempenhar atividades do dia a dia e a ter maior independência.”

A REVOLUÇÃO DAS PRÓTESES

Grande parte da verba para essas importantes descobertas vem de um projeto da Darpa chamado Revolutionizing Prosthetics [Revolucionando Próteses], com um fundo de 150 milhões de dólares que financia esses trabalhos desde 2006. Um dos propulsores do projeto é o coronel reformado do Exército dos Estados Unidos Geoffrey Ling, neurologista que serviu muitos anos no Iraque e no Afeganistão. Ele testemunhou horrorizado a carnificina produzida por minas terrestres nos campos de batalha. Em guerras anteriores, muitos desses corajosos soldados teriam morrido no local, mas hoje, com o uso de helicópteros e maior infraestrutura médica para remoção de feridos, muitos sobrevivem, apesar das graves lesões. Mais de 1.300 combatentes retornaram do Oriente Médio mutilados.

O dr. Ling se perguntou se haveria um meio científico de substituir esses membros perdidos. Financiado pelo Pentágono, ele pediu à sua equipe para apresentar soluções concretas num prazo de cinco anos. Esse pedido foi recebido com incredulidade. Ele recorda: “Acharam que estávamos loucos. Mas é na loucura que as coisas acontecem.”

Estimulada pelo entusiasmo ilimitado do dr. Ling, sua equipe fez milagres no laboratório. Por exemplo: o projeto Revolucionando Próteses financiou cientistas do Laboratório de Física Aplicada da Universidade Johns Hopkins, criadores do braço mecânico mais avançado da Terra, que faz quase todos os movimentos delicados de dedos, mãos e braços em três dimensões. A prótese tem os mesmos tamanho, força e agilidade de um braço humano. Embora seja feito de aço, se for coberto de plástico da cor da pele é quase indistinguível de um braço verdadeiro.

O braço foi ligado a Jan Sherman, uma tetraplégica com uma doença genética que destruiu a conexão entre o cérebro e o corpo, deixando-a completamente paralisada do pescoço para baixo. Na Universidade de Pittsburgh, foram colocados eletrodos diretamente na parte superior do cérebro dela, conectados a um computador e a um braço mecânico. Cinco meses após a cirurgia, Jan foi ao programa *60 Minutes*. Em rede nacional, ela usou o braço acenando alegremente para o público e cumprimentando o apresentador com um aperto de mão. Cumprimentou-o até com um soquinho de mão fechada para mostrar como o braço era sofisticado.

Ling diz: “Meu sonho é conseguirmos colocá-lo em todo tipo de paciente, com derrame, paralisia cerebral e idosos.”

Não só os cientistas, mas também os empresários estão de olho na interface cérebro-máquina (ICM). Querem incorporar permanentemente muitos desses inventos a seus planos de negócios. A ICM já penetrou no mercado jovem, na forma de videogames e brinquedos que usam sensores EEG para controlar objetos com a mente, tanto na realidade virtual como no mundo real. Em 2009, a NeuroSky comercializou o primeiro brinquedo com sensores EEG, o Mindflex, criado especificamente para mover uma bola através de um labirinto. Enquanto se usa o EEG Mindflex, a concentração aumenta a velocidade de um ventilador dentro do labirinto para impulsionar a bolinha pelos caminhos.

Os videogames controlados pela mente estão começando a estourar no mercado. Há 1.700 desenvolvedores de softwares trabalhando na NeuroSky, muitos deles dedicados aos acessórios Mindwave Mobile, com um orçamento de 129 milhões de dólares. Esses videogames vêm com um pequeno sensor de EEG portátil que é colocado em torno da testa e permite navegar na realidade virtual, onde os movimentos de um avatar são controlados mentalmente. Enquanto manobramos o avatar na tela, podemos usar armas de fogo, fugir do inimigo, passar de nível, marcar pontos etc., como em qualquer videogame comum, exceto que não precisamos usar as mãos.

“Haverá um ecossistema inteiro de novos jogadores, e a NeuroSky está muito bem posicionada para ser como a Intel dessa nova indústria”, afirma Alvaro Fernandez, da empresa de pesquisa de mercado SharpBrains. Além de disparar armas virtuais, o capacete de EEG também é capaz de acusar quedas de atenção. A NeuroSky vem sendo consultada por empresas preocupadas com acidentes de trabalhadores que perdem a concentração enquanto operam máquinas perigosas, ou que adormecem ao volante. Essa tecnologia pode salvar vidas, alertando o trabalhador ou o motorista quando estiver perdendo a concentração. O capacete de EEG pode tocar um alarme quando a pessoa começa a cochilar.

No Japão, o capacete já é moda em festas. Os sensores de EEG têm forma de orelhas de gato quando colocados na cabeça. As orelhas se levantam quando a atenção está concentrada e se abaixam quando dispersa. Nas festas, pode ser usado para indicar uma paquera somente com o pensamento, e assim as pessoas sabem se estão impressionando alguém.

Mas talvez a maior novidade na aplicação dessa tecnologia seja a pesquisa do dr. Miguel Nicolelis, da Duke University. Quando o entrevistei, ele disse que acha possível produzir muitos dispositivos encontrados somente na ficção científica.

Nicolelis mostrou que a interface cérebro-máquina é capaz de atravessar continentes. Ele colocou um macaco numa esteira rolante, com um chip no cérebro do animal conectado à internet. Em Kioto, no Japão, do outro lado do planeta, os sinais emitidos pelo macaco são usados para fazer um robô andar. Ao andar na esteira na Carolina do Norte, nos Estados Unidos, o macaco controla o robô no Japão, que executa os mesmos movimentos de caminhada. Usando apenas os sensores no cérebro e uma guloseima de recompensa, Nicolelis treinou esses macacos para controlar um robô humanoide chamado CB-1, no outro lado do mundo.

Agora ele está às voltas com um dos principais problemas da interface cérebro-máquina: a falta de sensibilidade. As próteses de mãos atuais não têm sensação tátil e, portanto, parecem estranhas ao corpo. Como não há tato, podem acidentalmente esmagar os dedos de alguém num aperto de mãos. Pegar uma casca de ovo com um braço mecânico é praticamente impossível.

Nicolelis espera contornar esse problema com uma interface direta de cérebro a cérebro. Nesse invento, as mensagens do cérebro são enviadas ao braço mecânico que, com seus sensores, devolve as mensagens diretamente para o cérebro, sem percorrer o tronco encefálico. Essa máquina de interface cérebro-máquina-cérebro (ICMC) possibilitaria um mecanismo de feedback livre, direto, permitindo a sensação de tato.

Nicolelis começou por ligar o córtex motor de macacos rhesus a braços mecânicos. Tais braços têm sensores que enviam sinais de volta ao cérebro por meio de eletrodos conectados ao córtex somatossensorial (que registra o sentido do tato). Os macacos recebiam uma guloseima após cada tentativa bem-sucedida, e aprendiam a usar esse aparato depois de quatro a nove tentativas.

Para conseguir isso, Nicolelis inventou um código para representar diferentes superfícies (lisas ou ásperas). Ele disse que “após um mês de treinamento essa parte do cérebro aprende o código artificial que criamos e começa a associá-lo às diferentes texturas. Isso é uma primeira demonstração de que podemos criar um canal sensorial”, que simula as sensações da pele.

Falei com ele que essa ideia parecia o “holodeck” de *Jornada nas estrelas*, em que os personagens passeiam num mundo virtual, mas quando esbarram em objetos virtuais os sentem como se fossem reais. Isso é chamado “tecnologia háptica”, que usa tecnologia digital para simular o sentido do tato. Nicolelis respondeu: “Sim, acho que é uma primeira demonstração de que algo como o holodeck possa ser criado num futuro próximo.”

O holodeck do futuro pode usar uma combinação das duas tecnologias. Primeiro, as pessoas no holodeck usam lentes de contato ligadas à internet de modo a enxergar um mundo virtual inteiramente novo para onde quer que olhem. O cenário nas lentes de contato muda instantaneamente ao toque de um botão. E se a pessoa tocar em qualquer objeto nesse mundo virtual, os sinais

enviados ao cérebro simulam a sensação de toque, usando a tecnologia de interface cérebro-máquina-cérebro. Desse modo, objetos no mundo virtual visto através das lentes de contato parecerão sólidos.

A interface cérebro-cérebro possibilitaria não só a tecnologia háptica, mas também uma “internet da mente”, ou uma mentenet, com contato direto de cérebro a cérebro. Em 2013, Nicolelis realizou algo saído diretamente de *Jornada nas estrelas*, uma “fusão mental” de dois cérebros. Ele começou com dois grupos de ratos, um na Duke University e outro em Natal, no Brasil. O primeiro grupo aprendeu a pressionar uma alavanca quando via uma luz vermelha. O segundo grupo aprendeu a pressionar a alavanca quando seu cérebro era estimulado por um sinal enviado por meio de um implante. A recompensa por apertar a barra era um gole de água. Depois, Nicolelis conectou os córtices motores dos cérebros dos dois grupos à internet por meio de um fio fino.

Quando o primeiro grupo de ratos viu a luz vermelha, foi enviado um sinal via internet para o segundo grupo, no Brasil, que então pressionou a alavanca. Em sete de dez tentativas, o segundo grupo de ratos respondeu aos sinais enviados pelo primeiro grupo. Foi a primeira demonstração de que sinais podem ser enviados e interpretados corretamente entre dois cérebros. Falta muito para a fusão mental da ficção científica, em que dois cérebros se fundem em um só, porque o processo ainda é primitivo e a amostra é pequena, mas essa já é uma prova do princípio de que uma mentenet é possível.

Em 2013, outro passo importante aconteceu quando cientistas foram além dos estudos com animais e demonstraram a primeira comunicação humana direta cérebro a cérebro, com um cérebro humano enviando mensagem a outro via internet.

Esse marco foi atingido na Universidade de Washington, com um cientista enviando um sinal cerebral (mova seu braço direito) a outro cientista. O primeiro cientista usou um capacete EEG e, jogando um videogame, disparou um canhão imaginando mover o braço direito, mas tomou o cuidado de não movê-lo fisicamente.

O sinal do capacete de EEG foi enviado pela internet para outro cientista, que estava usando um capacete magnético transcraniano colocado exatamente sobre a parte do cérebro que controla o braço direito. Quando o sinal chegou ao segundo cientista, seu capacete enviou um pulso magnético ao cérebro, o que fez seu braço direito se mover sozinho, involuntariamente. Assim, por controle remoto, um cérebro humano pôde controlar o movimento de outra pessoa.

Essa descoberta abre uma série de possibilidades, como trocar mensagens não verbais via internet. Talvez um dia você possa enviar a experiência de dançar um tango, fazer *bungee jumping* ou pular de paraquedas à sua lista de contatos de e-mail. Não apenas atividade física, mas também sentimentos e emoções poderão

ser enviados por comunicação cérebro a cérebro.

Nicolelis prevê o dia em que as pessoas do mundo inteiro poderão participar de redes sociais, não usando teclados, mas por meio da mente. Em vez de enviar e-mails, os usuários da mentenet poderão trocar pensamentos, emoções e ideias telepaticamente, em tempo real. Hoje um telefonema transmite apenas a informação da conversa e o tom de voz, nada mais. A videoconferência é um pouco melhor, pois permite ler a linguagem corporal da pessoa do outro lado. Mas a mentenet seria a grande revolução nas comunicações, tornando possível compartilhar numa conversa toda a informação mental, incluindo emoções, nuances e ressalvas. As mentes poderão trocar seus pensamentos e sentimentos mais íntimos.

ENTRETENIMENTO DE IMERSÃO TOTAL

O desenvolvimento da mentenet poderá também ter um impacto na indústria multibilionária do entretenimento. Nos anos 1920, a tecnologia para gravação de som e imagem em fita foi aperfeiçoada. Isso transformou a indústria do entretenimento com a transição do cinema mudo para o falado. Essa fórmula básica de combinar som e imagem não mudou muito no século passado. Mas no futuro tal indústria poderá fazer a próxima transição, gravando os cinco sentidos, incluindo cheiro, paladar e tato, além de toda a gama de emoções. Sondas telepáticas poderão manejar todos os sentidos e emoções que circulam no cérebro, produzindo uma imersão completa do público na história. Ao assistir a um filme de amor ou de suspense, nadaremos num mar de sensações, como se estivéssemos realmente lá, vivenciando as torrentes de sentimentos e emoções dos atores. Poderemos sentir o perfume da mocinha, o pavor da vítima num filme de terror e o prazer ao derrotar o bandido.

Essa imersão implica uma mudança radical no modo de fazer filmes. Primeiro, os atores vão precisar aprender a atuar com sensores de EEG/IRM e nanossondas gravando suas emoções e sensações. Isso trará um trabalho a mais para os atores, que precisarão atuar em cada cena simulando os cinco sentidos. É verdade que alguns atores não conseguiram fazer a transição do cinema mudo para o falado, mas talvez uma nova geração de atores possa atuar com os cinco sentidos. A edição exigirá não só cortar e colar o filme, mas combinar a gravação das várias sensações em cada cena. E por fim, o público, sentado confortavelmente, receberá todos esses sinais elétricos no cérebro. Em vez de óculos de 3D, os espectadores usarão algum tipo de sensor cerebral. As salas de cinema também terão que ser reequipadas para processar esses dados e enviá-los ao público.

CRIANDO A MENTENET

A criação de uma mentenet capaz de transmitir essas informações precisará ser feita em etapas. O primeiro passo é inserir nanossondas em partes importantes do cérebro, como o lobo temporal esquerdo, que controla a linguagem, e o lobo occipital, que controla a visão. Depois os computadores analisam e decodificam os sinais, e a informação é enviada pela internet, por cabos de fibra óptica.

Mais difícil seria inserir esses sinais no cérebro de outra pessoa, onde seriam processados por um receptor. Até o momento, o progresso nessa área se concentrou apenas no hipocampo, mas no futuro será possível inserir mensagens diretamente em outras partes do cérebro, correspondentes à sensibilidade ao som, luz, toque etc. Portanto, os cientistas ainda terão muito o que fazer para mapear os córtices cerebrais envolvidos nesses sentidos. Uma vez mapeados os córtices – como o hipocampo, que discutiremos no próximo capítulo –, será possível inserir em outro cérebro palavras, pensamentos, lembranças e experiências.

Nicolelis escreve que “não é inconcebível que nossa descendência humana consiga de fato reunir a técnica, a tecnologia e a ética necessárias para estabelecer uma mentenet funcional, um meio pelo qual bilhões de seres humanos estabeleçam, consensualmente, um contato direto e temporário com outros através apenas do pensamento. O que esse volume colossal de consciência coletiva poderá parecer, suscitar ou fazer, nem eu, nem ninguém hoje, é capaz de imaginar ou expressar”.

A MENTENET E A CIVILIZAÇÃO

Uma mentenet pode até mudar o curso da própria civilização. Toda vez que um novo sistema de comunicação foi introduzido provocou uma irrevogável aceleração de mudanças na sociedade, transportando-nos para a era seguinte. Na pré-história, durante milhares de anos nossos ancestrais foram nômades, andando em pequenas tribos, comunicando-se com os outros por linguagem corporal e grunhidos. A chegada da linguagem nos permitiu, pela primeira vez, comunicar símbolos e ideias complexas, o que facilitou o crescimento de vilas, e depois, de cidades. Nos últimos milhares de anos, a linguagem escrita nos permitiu acumular conhecimento e cultura através das gerações, possibilitando a ascensão da ciência, das artes, da arquitetura e de impérios imensos. O surgimento do telefone, rádio e TV estenderam o alcance da comunicação a continentes. Hoje a internet torna possível a ascensão de uma civilização planetária que ligará todos os continentes e pessoas do mundo inteiro. O próximo passo gigantesco será uma

mentenet planetária, em que o espectro completo dos sentidos, emoções, lembranças e pensamentos será compartilhado em escala global.

“SEREMOS PARTE DO SISTEMA OPERACIONAL”

Quando entrevistei Nicolelis, ele disse que se interessou por ciência muito cedo, ainda no Brasil, onde cresceu. Ele se lembra de ter visto o lançamento da Apollo à Lua, o que capturou a atenção do mundo. Para ele, foi uma proeza espantosa. E agora, disse ele, seu próprio “lançamento à Lua” é poder mover um objeto com a mente.

Ele começou a se interessar pelo cérebro ainda no ensino médio, onde achou um livro de Isaac Asimov intitulado *O cérebro humano*, de 1964. Mas ficou decepcionado no fim do livro. Não havia uma discussão sobre como todas as estruturas interagiam entre si para criar a mente (porque ninguém sabia a resposta na época). Esse momento mudou sua vida e ele soube que seu destino seria tentar entender os segredos do cérebro.

Cerca de dez anos atrás, ele me contou, começou a pensar seriamente em realizar a pesquisa de seu sonho de infância. Começou fazendo um ratinho controlar um dispositivo mecânico. “Colocamos no rato sensores que liam sinais elétricos emitidos pelo cérebro. Depois transmitimos esses sinais para uma pequena alavanca robótica que trazia água de uma fonte para a boca do rato. Assim o animal tinha que aprender a mover mentalmente o dispositivo robótico para beber a água. Foi a primeira demonstração de todos os tempos de que se pode conectar um animal a uma máquina de modo a conseguir operar a máquina sem mover o próprio corpo”, ele explicou.

Hoje ele consegue analisar não apenas 50, mas 1.000 neurônios no cérebro de um macaco, que podem reproduzir vários movimentos de diferentes partes do corpo. Assim o macaco consegue controlar diversos dispositivos, como braços mecânicos, e até imagens virtuais no ciberespaço. “Temos até um avatar macaco que pode ser controlado pelos pensamentos do animal, sem que precise fazer qualquer movimento”, ele disse. Isso é feito com o macaco assistindo a um vídeo em que há um avatar representando seu corpo. Depois, comandando mentalmente o movimento de seu corpo, o macaco faz o avatar se mover da maneira correspondente.

Nicolelis prevê o dia num futuro próximo em que jogaremos videogames e controlaremos computadores e outros aparelhos com a mente. “Seremos parte do sistema operacional dessas máquinas. Ficaremos imersos nelas por meio de mecanismos muito similares aos experimentos que estou descrevendo.”

EXOESQUELETOS

O próximo empreendimento do dr. Nicolelis é o Projeto Andar de Novo. Sua meta é nada menos que um exoesqueleto completo para o corpo, controlado pela mente. Nesse primeiro momento, o exoesqueleto parece ter saído de filmes do *Homem de ferro*. Na verdade, é um traje especial que encapsula o corpo inteiro, de maneira que os braços e as pernas podem ser movidos por meio de motores. Ele o chama de “robô vestível”. (Ver Figura 10.)

Seu objetivo, diz Nicolelis, é ajudar os paraplégicos a “andar com o pensamento”. Ele planeja usar tecnologia sem fio, “para que não fique nada espetado fora da cabeça. (...) Vamos alistar de vinte a trinta mil neurônios para comandarem uma vestimenta robótica de corpo inteiro. E, com o pensamento, o paciente poderá andar de novo e pegar objetos”.



Figura 10. O dr. Nicolelis espera que este exoesqueleto seja controlado pela mente de um tetraplégico.

Nicolelis sabe que precisa vencer uma série de obstáculos até o exoesqueleto se tornar realidade. Primeiro, é preciso criar uma nova geração de microchips que possam permanecer no cérebro com segurança durante anos. Segundo, é preciso criar sensores sem fio para o exoesqueleto andar sem empecilhos. Os sinais do cérebro são recebidos sem fio por um computador do tamanho de um

telefone celular, que provavelmente ficará preso em um cinto. Terceiro, serão necessários novos avanços para decifrar e interpretar sinais cerebrais via computador. Para os macacos, bastaram poucas centenas de neurônios para controlar os braços mecânicos. Para humanos, será preciso, no mínimo, muitos milhares de neurônios para controlar um braço ou uma perna. E quarto, é preciso encontrar uma fonte de energia que seja portátil e potente o bastante para ativar o exoesqueleto inteiro.

A meta de Nicolelis era muito ousada: ter o exoesqueleto pronto para a Copa do Mundo de 2014 no Brasil, onde um brasileiro tetraplégico deu o chute inicial. Ele me disse, com orgulho: “É a chegada brasileira à Lua.”

AVATARES E SUBSTITUTOS

No filme *Substitutos*, Bruce Willis faz o papel de um agente do FBI que investiga assassinatos misteriosos. Os cientistas criaram exoesqueletos com tanta perfeição que superaram as capacidades humanas. São criaturas mecânicas superfortes e com corpos perfeitos. Na verdade, são tão perfeitos que a humanidade passou a depender deles. As pessoas vivem o tempo todo em casulos, controlando mentalmente com tecnologia sem fio seus belíssimos androides substitutos. Em todos os lugares, há “gente” ocupada com afazeres, mas são apenas substitutos perfeitamente modelados. Seus usuários envelhecem convenientemente escondidos. A trama dá uma virada súbita, quando Bruce Willis descobre que a pessoa por trás de assassinatos pode estar ligada ao próprio cientista que inventou os substitutos. Isso o leva a questionar se os substitutos são uma bênção ou uma maldição.

E no grande sucesso *Avatar*, no ano de 2154 a maior parte dos minerais da Terra está esgotada, e uma companhia de mineração viaja para uma lua distante, chamada Pandora, no sistema estelar Alfa de Centauro, em busca de um metal raro, *unobtainium*. Essa lua é habitada por nativos chamados Na'vi, que vivem em harmonia com a incrível natureza do lugar. Para se comunicar com o povo de lá, trabalhadores especialmente treinados são colocados em casulos, onde aprendem a controlar com a mente o corpo geneticamente modificado de um nativo. Apesar da atmosfera perniciosa e do ambiente totalmente diferente da Terra, os avatares não têm dificuldade em viver naquele mundo. Porém, esse estranho relacionamento logo se deteriora quando a companhia mineradora descobre um rico depósito de *unobtainium* embaixo da árvore sagrada dos cerimoniais Na'vi. O conflito entre a mineradora, que quer destruir a árvore sagrada e escavar o solo para obter o metal raro, e os nativos, que veneram a árvore, é inevitável. Parece uma causa perdida para os nativos, até que um dos trabalhadores especialmente

treinados muda de lado e leva os Na'vi à vitória.

Avatares e substitutos são hoje elementos de ficção científica, mas um dia podem vir a ser uma ferramenta essencial da ciência. O corpo humano é frágil, talvez delicado demais para os rigores de muitas missões perigosas, inclusive viagens espaciais. Embora a ficção científica seja cheia de explorações heroicas de astronautas corajosos viajando até os confins da galáxia, a realidade é muito diferente. A radiação no espaço é tão intensa que os astronautas teriam que ser blindados, ou encarar o envelhecimento precoce, os efeitos nocivos da radiação, e até câncer. Explosões solares podem atingir uma nave espacial com radiação letal. Um simples voo transatlântico entre Estados Unidos e Europa expõe a pessoa a um milirém de radiação por hora, aproximadamente o mesmo que uma radiografia dental. No espaço sideral a radiação pode ser muitas vezes mais intensa, especialmente na presença de raios cósmicos e erupções solares. (Durante grandes tempestades solares, a Nasa de fato avisou aos astronautas na estação espacial para permanecerem em setores mais protegidos contra a radiação.)

Além disso, existem muitos outros perigos no espaço sideral, como micrometeoritos, efeitos de ausência de peso prolongada, e os problemas de adaptação a diferentes campos de gravidade. Após alguns meses de ausência de peso, o corpo perde uma grande porção de cálcio e sais minerais, deixando os astronautas incrivelmente fracos, mesmo com exercícios diários. Depois de um ano no espaço, os astronautas russos tiveram que sair das cápsulas se arrastando como minhocas. Além do mais, acredita-se que alguns efeitos da perda muscular e óssea sejam permanentes, e os astronautas sentirão pelo resto da vida as consequências da ausência de peso prolongada.

O perigo de meteoritos e campos de radiação intensa na Lua é tão grande que vários cientistas propuseram montar uma estação lunar permanente numa gigantesca caverna subterrânea para proteger os astronautas. Essas cavernas são formações naturais de tubos de lava, próximas a vulcões extintos. Mas o modo mais seguro de construir uma base lunar é com os astronautas sentados confortavelmente em casa. Assim ficam protegidos contra todos os riscos encontrados na Lua e, através de substitutos, podem cumprir as mesmas tarefas. Isso reduz consideravelmente os custos de viagens espaciais tripuladas, pois suprir as necessidades vitais de astronautas humanos custa muito caro.

Talvez, quando a primeira nave interplanetária alcançar um planeta distante, e um substituto de astronauta pisar em terreno alienígena, este poderá dar “um pequeno passo para a mente...”.

Um problema dessa abordagem é o tempo que leva para mensagens chegarem à Lua e além dela. Em pouco mais de um segundo uma mensagem de rádio viaja da Terra até a Lua, portanto substitutos na Lua podem ser facilmente controlados por astronautas na Terra. Mais difícil será se comunicar com

substitutos em Marte, já que leva vinte minutos ou mais para os sinais de rádio chegarem ao Planeta Vermelho.

Mas a utilização de substitutos tem implicações práticas mais perto daqui. No Japão, o acidente com o reator de Fukushima em 2011 resultou num prejuízo de bilhões de dólares. Como os trabalhadores só podem entrar nas áreas com níveis letais de radiação por poucos minutos, a limpeza final pode levar até quarenta anos. Infelizmente, não existem robôs suficientemente avançados para suportar a intensidade dessa radiação e fazer os reparos necessários. Os únicos robôs usados em Fukushima são bastante primitivos, basicamente câmeras simples acopladas no alto de um computador sobre rodas. Um verdadeiro autômato, que possa pensar por si mesmo (ou ser controlado por um operador remoto) e fazer reparos em campos de alta radiação está a décadas de distância no futuro.

A falta de robôs industriais foi um problema crítico também para os soviéticos no acidente de 1986 em Chernobyl, na Ucrânia. Os homens enviados ao local do acidente para apagar as chamas tiveram uma morte horrível, devido à radiação. Mikhail Gorbachev então ordenou que a força aérea apagasse o incêndio “a seco”, despejando de helicóptero cinco mil toneladas de uma mistura de cimento e areia rica em boro. Os níveis de radiação eram tão altos que foram recrutados 250 mil homens para conter a catástrofe. Cada um deles só podia passar alguns minutos no prédio do reator fazendo os reparos. Muitos receberam a dose máxima de radiação permitida para a vida toda. Cada um ganhou uma medalha. Essa imensa operação foi o maior feito já realizado pela engenharia civil. Não poderia ter sido realizado pelos robôs atuais.

É verdade que a Honda Corporation construiu um robô que possivelmente poderá entrar em ambientes mortalmente radiativos, mas ainda não está pronto. Os cientistas da Honda colocaram na cabeça de um operário um sensor de EEG, conectado a um computador que analisa as ondas cerebrais. O computador é conectado a um rádio que envia mensagens ao robô, chamado Asimo (Advanced Step in Innovative Mobility). Desse modo, alterando suas próprias ondas cerebrais, o homem pode controlar o Asimo puramente com o pensamento.

Infelizmente, esse robô seria incapaz de fazer reparos em Fukushima, pois só consegue executar quatro movimentos básicos (e todos se limitam a mover a cabeça e os ombros), e centenas de movimentos são necessários para fazer reparos numa usina nuclear devastada. O sistema ainda não está desenvolvido o suficiente nem para efetuar tarefas simples, como girar uma chave de fenda ou bater um martelo.

Outros grupos também exploraram a possibilidade de criar robôs controlados pela mente. Na Universidade de Washington, o dr. Rajesh Rao criou um robô similar, controlado por uma pessoa usando um capacete de EEG. Esse reluzente robô humanoide tem 60 centímetros de altura e se chama Morpheus (como o personagem do filme *Matrix* e o deus grego dos sonhos). Um aluno põe o

capacete de EEG e faz gestos, como mover a mão, gerando um sinal de EEG que é gravado por um computador. O computador vai criando uma biblioteca de sinais de EEG correspondentes aos movimentos específicos de cada membro. Depois o robô é programado para mover a mão cada vez que o sinal correspondente de EEG for enviado a ele. Assim, quando a pessoa pensa em mover a mão, o Morpheus move a mão também. Quando alguém coloca o capacete de EEG pela primeira vez, o computador leva cerca de dez minutos para calibrar os sinais do cérebro. Depois, domina-se a técnica de fazer gestos mentalmente para controlar o robô. Por exemplo: você pode fazer o robô andar em sua direção, pegar um bloco em uma mesa, caminhar dois metros até outra mesa e pôr o bloco lá.

A pesquisa está progredindo rapidamente também na Europa. Em 2012, cientistas da École Polytechnique Fédérale de Lausanne, na Suíça, revelaram sua última conquista: um robô controlado telepaticamente por sensores de EEG, cujo controlador fica a 100 quilômetros de distância. O robô tem a aparência de um aspirador de pó robótico Roomba, já encontrado em muitos lares. Mas, na verdade, é altamente sofisticado, equipado com uma câmera que consegue achar caminho num escritório repleto de mesas. Um paciente paralisado pode, por exemplo, olhar para uma tela de computador conectada a uma câmera de vídeo no robô a quilômetros de distância, e ver pelos olhos do robô. Com o pensamento, o paciente é capaz de controlar os movimentos da máquina para contornar obstáculos.

No futuro, podemos imaginar os trabalhos mais perigosos sendo feitos por robôs controlados por humanos dessa maneira. O dr. Nicoletti diz: “Provavelmente seremos capazes de operar por controle remoto embaixadores e emissários, robôs e aeronaves de várias formas e tamanhos, enviados como nossos representantes para explorar planetas e estrelas em lugares distantes no universo.”

Por exemplo: em 2010 o mundo viu, com horror, cinco milhões de barris de petróleo bruto se espalhando pelo golfo do México. O vazamento na plataforma Deepwater Horizon foi um dos maiores desastres com petróleo da história, e os engenheiros se viram praticamente impotentes durante três meses. Submarinos robóticos, guiados por controle remoto, se atrapalharam durante semanas tentando tampar o poço, porque não tinham a destreza e versatilidade necessárias para realizar a missão embaixo da água. Se pudessem usar submarinos substitutos, muito mais sensíveis para manipular ferramentas, o vazamento poderia ter sido controlado nos primeiros dias, evitando um prejuízo de bilhões em danos e ações judiciais.

Outra possibilidade é que minúsculos submarinos substitutos possam um dia entrar no corpo humano e realizar cirurgias delicadas lá dentro. Essa ideia foi explorada no filme *Viagem fantástica*, estrelado por Raquel Welch, em que um

submarino é encolhido até o tamanho de uma célula de sangue e injetado na corrente sanguínea de uma pessoa com um coágulo no cérebro. O encolhimento de átomos viola as leis da física quântica, mas um dia será possível injetar MEMS (da sigla em inglês para sistemas microeletromecânicos) do tamanho de células na corrente sanguínea humana. Os MEMS são máquinas incrivelmente minúsculas, que cabem na ponta de um alfinete. Os MEMS usam a mesma tecnologia empregada no Vale do Silício, que possibilita colocar milhões de transistores numa pastilha do tamanho de uma unha. Uma máquina completa com engrenagens, alavancas, polias, e até motores pode ser menor do que o ponto final dessa frase. Um dia, será possível colocar um capacete telepático e comandar um submarino MEMS usando tecnologia sem fio para fazer uma cirurgia dentro de um paciente.

Assim, a tecnologia MEMS pode abrir um campo inteiramente novo na medicina, criando máquinas microscópicas para entrar no corpo. Os submarinos MEMS poderiam até guiar nanossondas dentro do cérebro para conectar neurônios específicos. Dessa maneira, as nanossondas seriam capazes de receber e transmitir sinais do grupo de neurônios envolvidos em determinados comportamentos. A técnica eliminaria o procedimento pouco preciso de inserção de eletrodos no cérebro.

O FUTURO

Em curto prazo, todos esses avanços notáveis conquistados nos laboratórios do mundo inteiro podem aliviar o martírio dos que sofrem de paralisias e outras deficiências. Usando o poder da mente, eles serão capazes de se comunicar com as pessoas queridas, controlar cadeiras de rodas e camas, andar com membros mecânicos guiados mentalmente, manejar utensílios domésticos e levar uma vida seminormal.

Mas, em longo prazo, esses avanços poderão ter profundas implicações práticas e econômicas no mundo. Em meados deste século, interagir mentalmente com computadores pode ser comum. Como a informática é uma indústria de muitos trilhões de dólares, capaz de produzir jovens bilionários e corporações da noite para o dia, avanços na interface mente-computador irão reverberar em Wall Street – e também na nossa casa.

Todos os acessórios que usamos para acessar computadores (mouse, teclado etc.) podem acabar desaparecendo. No futuro, poderemos simplesmente dar comandos mentais e nossos desejos serão cumpridos por minúsculos chips ocultos no ambiente. Sentados no escritório, passeando no parque, vendo vitrines, ou parados, só relaxando, nossa mente pode estar interagindo com um monte de

chips escondidos, permitindo-nos saber nosso saldo no banco, fazer reservas num restaurante ou comprar ingressos para o teatro.

Os artistas também podem usufruir bastante dessa tecnologia. Ao visualizarem uma obra, a imagem poderá ser enviada via sensores de EEG para uma tela holográfica em 3D. Como a imagem mental não é tão exata quanto o objeto original, o artista pode aprimorar depois a imagem em 3D e visualizar a obra retocada. Após alguns ciclos, ele pode produzir a imagem final numa impressora 3D.

Da mesma forma, engenheiros poderão criar modelos em escala de pontes, túneis, aeroportos, usando apenas a imaginação. Poderão alterar as plantas rapidamente, com o pensamento. Peças de máquinas sairão da tela do computador para uma impressora 3D.

Alguns críticos, porém, alegaram que esses poderes telecinéticos têm uma grande limitação: a falta de energia. No cinema, seres superiores têm o poder de mover montanhas com o pensamento. No filme *X-Men: O confronto final*, o supervilão Magneto tem a capacidade de mover a ponte Golden Gate simplesmente apontando os dedos, mas o corpo humano só consegue reunir, em média, cerca de um quinto de cavalo de força, o que é muito pouco para realizar as proezas que vemos nas histórias em quadrinhos. Portanto, todos os feitos hercúleos de seres superiores telecinéticos parecem ser pura fantasia.

Mas há uma solução para o problema da energia. Será possível conectar o pensamento a uma fonte de energia que amplia milhões de vezes a nossa força. Assim, conseguiremos chegar perto de ter a força de um deus. Num episódio de *Jornada nas estrelas*, a tripulação viaja para um planeta distante onde encontra uma criatura divina que diz ser Apolo, o deus grego do Sol. Ele é capaz de realizar mágicas que deslumbram a tripulação. E diz até que tinha ido à Terra muitas eras antes, onde os terráqueos o adoravam. Mas a tripulação, que não acredita em deuses, suspeita de charlatanice. Depois descobrem que o “deus” sabia controlar mentalmente uma fonte de força escondida, que fazia todas as mágicas. Quando a fonte é destruída, ele se torna um mero mortal.

Da mesma forma, no futuro nossa mente talvez controle uma fonte de energia que nos dê superpoderes. Por exemplo: um mestre de obras poderá usar telepaticamente uma fonte de força para energizar máquinas pesadas, e construir sozinho casas e edifícios complexos, usando apenas o poder da mente. Todos os pesos seriam levantados pela força da fonte de energia, e o construtor seria como um maestro, orquestrando todo o movimento de guindastes gigantes e escavadeiras potentes, por meio do pensamento.

A ciência está começando a se equiparar à ficção científica em outro aspecto. A saga de *Star Wars* se passa numa época em que civilizações se espalharam por toda a galáxia. A paz é mantida pelos Cavaleiros de Jedi, um grupo de guerreiros altamente qualificados para usar o poder da “Força”, ler pensamentos e manejar

sabres de luz.

Mas não precisamos esperar até que toda a galáxia esteja colonizada para começarmos a pensar na Força. Como vimos, alguns aspectos da Força já são possíveis hoje, como penetrar no pensamento de outros por meio de eletrodos de ECOG e capacetes de EEG. Os poderes telecinéticos dos Cavaleiros de Jedi também serão possíveis quando aprendermos a acionar uma fonte de força com a mente. Os Cavaleiros de Jedi, por exemplo, fazem surgir um sabre de luz a um simples gesto com as mãos, mas já podemos fazer a mesma proeza explorando o poder do magnetismo (semelhante ao ímã do aparelho de IRM, capaz de arremessar um martelo para o outro lado da sala). Ao ativar mentalmente a fonte de força, podemos pegar um sabre de luz do outro lado da sala, usando a tecnologia de hoje.

O PODER DE UM DEUS

A telecinesia é um poder geralmente reservado a uma divindade ou a um super-herói. No universo dos super-heróis dos filmes de Hollywood, talvez a personagem mais poderosa seja Fênix, a mulher telecinética que move objetos à vontade. Como membro dos X-Men, ela levanta máquinas pesadas, aviões a jato, estanca inundações, tudo com o poder da mente. No entanto, quando é finalmente consumida pelo lado negro de seu poder, ela entra numa fúria cósmica capaz de incinerar sistemas solares inteiros e destruir estrelas. Seu poder é tão grande e incontrolável que acaba por levá-la à destruição.

Mas até onde a ciência pode chegar no domínio dos poderes telecinéticos?

No futuro, mesmo tendo uma fonte de força externa para amplificar nossos pensamentos, é improvável que pessoas com poderes telecinéticos sejam capazes de mover mentalmente objetos simples, como um lápis ou uma xícara de café. Como mencionamos, existem apenas quatro forças conhecidas que regem o universo, e nenhuma delas consegue mover objetos sem uma fonte de força externa. O magnetismo chega perto, mas pode mover apenas objetos magnéticos. Objetos de plástico, água ou madeira atravessam facilmente campos magnéticos. A levitação, truque muito usado em shows de mágica, está além de nossa competência científica.

Assim, mesmo com fornecimento de força externa, é improvável que alguém possa mover objetos à vontade por telecinesia. Existe, porém, uma tecnologia que se aproxima disso, envolvendo a capacidade de transformar um objeto em outro.

Essa tecnologia é chamada “matéria programável”, e foi tema de intensa pesquisa para a Intel. A ideia por trás da matéria programável é criar objetos

feitos de minúsculos “cátomos”,^[2] que são chips microscópicos de computador. Cada cátomo pode ser controlado sem fio. Pode ser programado para mudar a carga elétrica em sua superfície a fim de se ligar a outros cátomos de diversas maneiras. Tendo a carga elétrica programada de certa maneira, os cátomos se ligariam para formar, digamos, um telefone celular. Apertando-se um botão para mudar sua programação, os cátomos se rearranhariam para formar outro objeto, como um laptop, por exemplo.

Vi uma demonstração dessa tecnologia na Universidade Carnegie Mellon, em Pittsburgh, onde cientistas conseguiram criar um chip do tamanho de uma cabeça de alfinete. Para examinar os cátomos, tive que entrar numa “sala limpa”, vestindo um macacão branco especial, botas de plástico e touca, para evitar que a menor partícula de pó entrasse na sala. Então vi, pelo microscópio, o intrincado circuito dentro de cada cátomo, que torna possível programar, sem fio, uma mudança de carga elétrica em sua superfície. Da mesma maneira que podemos programar softwares hoje, no futuro pode ser possível programar hardwares.

O próximo passo é determinar se os cátomos podem se combinar para formar objetos úteis, e se podem ser alterados ou tomar a forma de outro objeto conforme nossa vontade. Pode levar até meio século para conseguirmos utilizar protótipos de matéria programável. Dada a complexidade de programar bilhões de cátomos, deve ser criado um computador especial para orquestrar as cargas de todos eles. Talvez no fim deste século seja possível controlar mentalmente esse computador para podermos transformar um objeto em outro. Não precisaremos memorizar as cargas e a configuração interna de um objeto. Basta dar um comando mental para o computador transformar um objeto em outro.

Depois teremos catálogos dos vários objetos programáveis, como móveis, utensílios e aparelhos eletrônicos. Por comunicação telepática com o computador, será possível transformar objetos. Redecorar a sala, reformar a cozinha, comprar presentes de Natal, tudo poderá ser feito mentalmente.

MORAL DA HISTÓRIA

Realizar todos os desejos é algo exclusivo das divindades. Entretanto, há também um lado negativo nesse poder celestial. Todas as tecnologias podem ser usadas para o bem e para o mal. Em última análise, a ciência é uma faca de dois gumes. Um dos lados pode acabar com a pobreza, a doença e a ignorância. O outro lado pode acabar com as pessoas, de várias maneiras.

Essas tecnologias podem tornar as guerras cada vez mais cruéis. Talvez um dia o combate corpo a corpo seja entre substitutos armados com um arsenal de alta

tecnologia. Os soldados verdadeiros, sentados em segurança a milhares de quilômetros de distância, podem disparar rajadas das mais avançadas armas, pouco se importando com os danos colaterais infligidos aos civis. As guerras travadas com substitutos podem preservar a vida dos próprios soldados, mas também causar danos horrendos a civis e a propriedades.

O maior problema é que esse poder talvez seja grande demais para ser controlado por simples mortais. No romance *Carrie, a estranha*, Stephen King explora o mundo de uma garota sempre infernizada pelos colegas. Era rejeitada pela turma e sua vida se tornou uma série infundável de insultos e humilhações. Mas seus perseguidores não sabiam de um detalhe: ela era telecinética.

Depois de suportar vários insultos e levar um banho de sangue com seu vestido novo na festa de formatura, ela finalmente explode. Reunindo todo seu poder telecinético, Carrie prende os colegas e os elimina, um a um. Num gesto final, ela decide queimar a escola inteira. Mas seu poder telecinético era grande demais para ser controlado e Carrie acaba morrendo no incêndio que ela mesma provocou.

O enorme poder da telecinesia, além do perigo de ser um tiro pela culatra, tem outro problema. Ainda que se tomem todas as precauções para entender e controlar esse poder, ele pode nos destruir se, ironicamente, obedecermos a todos os nossos pensamentos e comandos. Assim, os pensamentos que nós mesmos produzimos podem nos destruir.

O filme *Planeta proibido*, de 1956, é baseado numa peça de William Shakespeare, *A tempestade*, que começa com um feiticeiro e a filha isolados numa ilha deserta. Mas em *Planeta proibido*, o professor e a filha vivem isolados num planeta distante que um dia foi habitado pelos krell, uma civilização milhões de anos mais avançada que a nossa. Sua maior realização foi um invento que lhes deu o poder da telecinesia, o poder de controlar mentalmente a matéria em todas as suas formas. Tudo o que desejavam se materializava imediatamente diante deles. Tinham o poder de remodelar a própria realidade como bem entendessem.

Mas às vésperas do grande triunfo, quando fariam o invento funcionar, os krell desapareceram sem deixar vestígios. O que teria destruído essa civilização tão avançada?

Quando um grupo de humanos aterrissou no planeta para resgatar o professor e a filha, descobriram que havia um terrível monstro assolando o planeta, matando membros da tripulação. Por fim, um dos tripulantes descobre o segredo por trás do monstro e dos krell. Antes de morrer, ele balbucia “Monstros do id”.

A terrível verdade subitamente se revela para o professor. Na noite em que os krell acionaram a máquina de telecinesia, eles adormeceram. Todos os seus desejos recalçados no id se materializaram. Submersos no subconsciente daquelas criaturas altamente desenvolvidas estavam os desejos e impulsos

animalescos de seu passado milenar. Todas as fantasias, todos os desejos de vingança se realizaram de repente e assim essa grande civilização destruiu a si mesma da noite para o dia. Havia conquistado tantos mundos, mas havia algo que não puderam controlar: sua própria mente inconsciente.

É uma lição para quem deseja liberar o poder da mente. Na mente se encontram as mais nobres conquistas e pensamentos da humanidade. Mas lá também se encontram os monstros do id.

MUDANDO QUEM SOMOS: A MEMÓRIA E A INTELIGÊNCIA

Até aqui discutimos o poder da ciência para ampliar nossas capacidades mentais via telepatia e telecinesia. Nós, basicamente, permanecemos os mesmos. Esses adventos nada fazem para mudar a essência de quem somos. Contudo, estão abrindo uma nova fronteira que altera a própria natureza do que significa ser humano. As novidades da genética, do eletromagnetismo e das terapias medicamentosas podem, num futuro próximo, alterar nossa memória e até mesmo expandir nossa inteligência. A ideia de recuperar uma lembrança, aprender técnicas complexas da noite para o dia e ficar superinteligente vem lentamente se desligando do mundo da ficção científica.

Sem a memória estamos perdidos, à deriva num mar inútil de estímulos sem objetivo, incapazes de entender o passado ou a nós mesmos. Então, o que acontecerá se algum dia pudermos inserir em nosso cérebro memórias artificiais? O que acontecerá quando pudermos nos tornar mestres em qualquer disciplina simplesmente baixando um arquivo em nossa memória? E o que acontecerá se não soubermos mais a diferença entre lembranças verdadeiras e falsas? Então, quem seremos?

Os cientistas estão deixando de ser observadores passivos da natureza e passando a modelar ativamente a natureza. Isso significa que poderemos ser capazes de manipular memória, pensamento, inteligência e consciência. Em vez de simplesmente testemunhar os intrincados mecanismos da mente, no futuro será possível orquestrá-los.

Então, passemos agora para a seguinte pergunta: Será possível fazer download de memórias?

2. Ou átomos claytrônicos, em português. A claytrônica é uma nova área da engenharia mecatrônica que trabalha com nanorrobôs reconfiguráveis. (N. do. R.T.)

Se nosso cérebro fosse simples a ponto de ser entendido, não seríamos inteligentes a ponto de entendê-lo.

– ANÔNIMO

5 MEMÓRIAS E PENSAMENTOS FEITOS SOB MEDIDA

Neo é “O Escolhido”. Somente ele pode levar a humanidade derrotada à vitória contra as Máquinas. Somente Neo pode destruir a Matrix, que implantou falsas lembranças no nosso cérebro para nos controlar.

Numa cena já clássica do filme *Matrix*, as Sentinelas do Mal, que guardam a Matrix, finalmente encurralam Neo. Parece que a última esperança da humanidade está prestes a terminar. Mas Neo já tinha inserido um eletrodo na nuca que podia transferir instantaneamente toda a habilidade das artes marciais para dentro de seu cérebro. Em segundos, ele é um mestre de caratê capaz de vencer as Sentinelas com chutes voadores fantásticos e golpes bem aplicados.

Em *Matrix*, aprender a extraordinária arte de um carateca faixa preta é tão fácil quanto inserir um eletrodo no cérebro e fazer um download. Talvez um dia nós também possamos importar memórias, o que aumentará bastante nossas habilidades.

Mas o que acontece quando as lembranças transferidas para o cérebro são falsas? No filme *O vingador do futuro*, Arnold Schwarzenegger tem lembranças falsas plantadas no cérebro, de modo que a distinção entre realidade e ficção é totalmente difusa. Ele luta valentemente contra os malvados em Marte até o fim do filme, quando se dá conta de que ele próprio é o líder dos malvados. Fica chocado ao descobrir que suas lembranças de ser um cidadão normal, respeitador das leis, são completamente fabricadas.

Hollywood adora fazer filmes que exploram o fascinante, porém ficcional, mundo da memória artificial. É claro que tudo isso é impossível com a tecnologia atual, mas podemos sonhar com um dia, daqui a algumas décadas, em que a memória artificial poderá de fato ser inserida no cérebro.

COMO LEMBRAMOS

Como na história de Phineas Gage, o estranho caso de Henry Gustav Molaison, conhecido na literatura científica como HM, causou tanto furor no campo da neurologia que levou a avanços fundamentais para o entendimento do hipocampo na formulação da memória.

Aos 9 anos, HM sofreu um acidente, resultando em ferimentos na cabeça que geraram convulsões graves. Em 1953, quando ele tinha 25 anos, submeteu-se a uma operação que aliviou totalmente os sintomas. Mas outro problema surgiu porque os cirurgiões cometeram um erro e retiraram parte do hipocampo dele. A princípio, HM parecia normal, mas logo ficou aparente que algo estava errado. Ele não conseguia reter a memória recente. Vivía todo o tempo no presente,

cumprimentando as mesmas pessoas várias vezes por dia, com as mesmas expressões, como se as estivesse vendo pela primeira vez. Tudo que entrava na memória durava poucos minutos e desaparecia. Como Bill Murray em *Feitiço do tempo*, HM ficou condenado a reviver o mesmo dia inúmeras vezes, pelo resto da vida. Mas ao contrário do personagem de Bill Murray, HM era incapaz de lembrar das últimas repetições. Sua memória de longo prazo, porém, ficou relativamente intacta, e ele se lembrava da vida antes da cirurgia. Mas, sem o funcionamento do hipocampo, HM era incapaz de gravar novas experiências. Por exemplo: ele ficava horrorizado quando se olhava no espelho porque via o rosto de um velho, e pensava que ainda tinha 25 anos. Mas felizmente a lembrança de ficar horrorizado também logo desaparecia nas brumas do esquecimento. Em certo sentido, HM era como um animal com Nível II de consciência, incapaz de recordar o passado imediato ou simular o futuro. Sem o funcionamento do hipocampo, ele regrediu do Nível III para o Nível II de consciência.

Hoje, os avanços da neurociência nos dão uma visão mais clara de como as lembranças são formadas, guardadas e revividas. “As peças só se encaixaram nos últimos anos, em consequência de dois desenvolvimentos técnicos – os computadores e a varredura cerebral moderna”, diz o dr. Stephen Kosslyn, neurologista de Harvard.

Como sabemos, as informações sensoriais (por exemplo, visão, tato, paladar) passam primeiro pelo tronco cerebral e pelo tálamo, que age como uma estação de retransmissão, direcionando os sinais para os vários lobos sensoriais do cérebro, onde são avaliados. A informação processada chega ao córtex pré-frontal, onde entra na consciência e forma o que chamamos de memória de curto prazo, abrangendo de alguns segundos a minutos. (Ver Figura 11.)

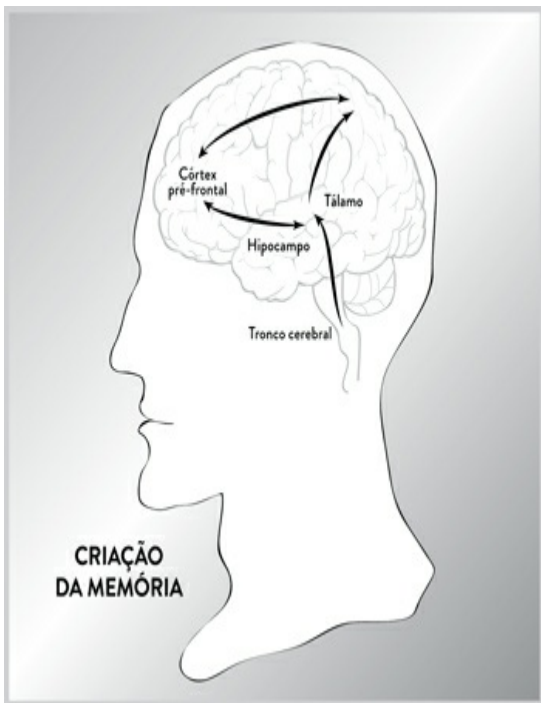


Figura 11. Aqui vemos o trajeto da criação de lembranças. Os impulsos dos sentidos passam pelo tronco cerebral, pelo tálamo e os vários córtices, até o córtex pré-frontal. Daí passam para o hipocampo, onde formam a memória de longo prazo.

Para que a memória tenha longa duração, a informação precisa passar pelo

hipocampo, onde as lembranças são divididas em várias categorias. Em vez de guardar todas as lembranças numa só área do cérebro, como um gravador de fita ou um disco rígido, o hipocampo redireciona fragmentos a vários córtices. O armazenamento de memória é assim muito mais eficaz do que se fosse sequencial. Se as lembranças humanas fossem armazenadas sequencialmente, como num disco rígido de computador, seria preciso uma área enorme para a memória. Na verdade, no futuro, até os sistemas de armazenamento digital podem copiar esse procedimento do cérebro vivo, em vez de armazenar lembranças inteiras sequencialmente. Por exemplo: as lembranças emocionais são armazenadas na amígdala, e as palavras são gravadas no lobo temporal. As cores e outras informações visuais são coletadas no lobo occipital, e a sensação de tato e movimento residem no lobo parietal. Até o momento, os cientistas identificaram mais de vinte categorias de memórias estocadas em diferentes partes do cérebro, incluindo frutas e legumes, plantas, animais, partes do corpo, cores, números, letras, substantivos, verbos, nomes próprios, rostos, expressões faciais, além de várias emoções e sons.

Uma única lembrança – um passeio no parque, por exemplo – contém informações que são fragmentadas e estocadas em várias regiões do cérebro. Mas trazer de volta um único aspecto dessa lembrança (por exemplo, o cheiro de grama recém-cortada) pode subitamente pôr o cérebro em ação para juntar os outros fragmentos e formar uma recordação coesa. O objetivo final da pesquisa sobre memória é saber como esses fragmentos espalhados são reunidos quando revivemos uma experiência. Isso se chama “problema da ligação”, e a solução deve explicar muitos aspectos enigmáticos da memória. Por exemplo: o dr. Antonio Damasio analisou pacientes que sofreram derrame e ficaram incapazes de identificar apenas uma categoria, mesmo conseguindo recordar do resto. Isso porque o derrame afetou somente uma determinada área do cérebro, onde aquela categoria estava armazenada.

O problema da ligação é ainda mais complicado porque as nossas lembranças e experiências são altamente pessoais. As lembranças são individualizadas, de modo que as categorias de memória não são as mesmas para todos. Provedores de vinho, por exemplo, têm várias categorias para rotular variações sutis de sabor, enquanto os físicos têm categorias diferentes para certas equações. Afinal, as categorias são subprodutos da experiência, e, portanto, pessoas diferentes têm categorias diferentes.

Uma solução recente para o problema da ligação se utiliza do fato de que há vibrações eletromagnéticas oscilando por todo o cérebro a aproximadamente 40 ciclos por segundo e podem ser captadas por EEG. Um fragmento de memória pode vibrar em uma frequência muito precisa e estimular outro fragmento de memória guardado numa parte distante do cérebro. Antes, pensava-se que as lembranças seriam armazenadas fisicamente juntas, mas essa nova teoria diz

que as lembranças não são unidas espacialmente, e, sim, temporalmente, por vibrações em uníssono. Se essa teoria for verdadeira, podemos dizer que há vibrações eletromagnéticas fluindo continuamente em todo o cérebro, ligando diferentes regiões e, assim, recriando lembranças inteiras. Consequentemente, o constante fluxo de informação entre o hipocampo, o córtex pré-frontal, o tálamo e os outros córtices pode não ser completamente neural. Um pouco desse fluxo pode estar na forma de ressonâncias através das diferentes estruturas cerebrais.

GRAVANDO UMA LEMBRANÇA

Infelizmente, HM morreu em 2008, aos 82 anos, antes de poder se beneficiar de alguns resultados sensacionais alcançados pela ciência: a capacidade de criar um hipocampo artificial e inserir lembranças no cérebro. Isso é algo que vem diretamente da ficção científica, mas cientistas da Universidade Wake Forest e da Universidade do Sul da Califórnia fizeram história em 2011, quando conseguiram gravar uma lembrança surgida em ratos e a armazenaram num computador. Foi um experimento de prova de conceito, mostrando que o sonho de importar memória para o cérebro pode um dia se tornar realidade.

A princípio, a própria ideia de fazer o download de lembranças no cérebro parece um sonho impossível porque, como vimos, elas são criadas pelo processamento de uma variedade de experiências sensoriais, que em seguida são armazenadas em múltiplos lugares no neocórtex e no sistema límbico. Mas sabemos, a partir do caso de HM, que existe um lugar por onde todas as lembranças fluem e são convertidas em memória de longo prazo: o hipocampo. O pesquisador Theodore Berger, da USC, diz: “Se não for feito no hipocampo, não poderá ser feito em lugar nenhum.”

Os cientistas da Wake Forest e da USC partiram da observação, com base em varreduras cerebrais, de que há pelo menos dois conjuntos de neurônios no hipocampo de um rato, chamados CA1 e CA3, que se comunicam entre si cada vez que o rato aprende uma tarefa. Depois de treinar ratinhos para apertar duas barras, uma após a outra, a fim de obter água, os cientistas revisaram as descobertas e tentaram decodificar as mensagens, o que foi frustrante a princípio, pois os sinais entre os dois conjuntos de neurônios não pareciam seguir um padrão. Mas, ao monitorar os sinais milhões de vezes, acabaram conseguindo determinar qual impulso elétrico produzia qual resultado. Usando sondas no hipocampo dos ratos, os cientistas puderam gravar os sinais entre CA1 e CA3 quando os ratos aprenderam a usar as duas barras em sequência.

Depois os cientistas injetaram nos ratos uma substância química especial, fazendo-os esquecer a tarefa. Por fim, reproduziram a lembrança no cérebro dos

mesmos ratos. De forma notável, a lembrança daquela tarefa retornou e os ratinhos conseguiram reproduzi-la como da primeira vez. Em essência, eles tinham criado um hipocampo artificial, com a capacidade de duplicar a memória digital. “Liga o botão, o animal tem memória; desliga o botão, e ele não tem mais”, diz o dr. Berger. “É um passo muito importante porque é a primeira vez que encaixamos todas as peças.”

Como disse Joel Davis, da Marinha dos EUA, que patrocinou a pesquisa, “já foi dada a partida para o uso de implantes que aumentem a competência. É só uma questão de tempo”.

Com tanta coisa em jogo, não é de surpreender que essa área de pesquisa esteja se desenvolvendo tão rapidamente. Em 2013, foi feita outra descoberta, dessa vez no MIT, por cientistas que implantaram não só lembranças comuns num rato, mas também lembranças falsas. Isso significa que, um dia, lembranças de eventos nunca acontecidos poderão ser implantadas no cérebro, o que terá um impacto profundo em áreas como educação e entretenimento.

Os cientistas do MIT usaram uma técnica chamada optogenética (que discutiremos mais no capítulo 8). Tal técnica permite focalizar uma luz em neurônios específicos para ativá-los. Usando esse método notável, os cientistas conseguem identificar os neurônios específicos responsáveis por certas lembranças.

Digamos que um rato entre em algum lugar e leve um choque. Os neurônios responsáveis pela lembrança desse evento doloroso podem ser isolados e registrados por meio da análise do hipocampo. Depois o rato é colocado em outro lugar, onde não há perigo algum. Acendendo luz em uma fibra óptica, pode-se usar a optogenética para ativar a lembrança do choque, e então o rato tem uma reação de medo, embora esse segundo lugar seja totalmente seguro.

Assim os cientistas do MIT conseguiram não só implantar uma lembrança comum, mas também a lembrança de algo que não aconteceu. Um dia, essa técnica poderá dar aos educadores a capacidade de implantar na memória habilidades específicas para requalificar trabalhadores, ou dar a Hollywood formas inteiramente novas de entretenimento.

HIPOCAMPO ARTIFICIAL

Atualmente, o hipocampo artificial é primitivo, capaz de gravar apenas uma lembrança de cada vez. Mas os cientistas planejam aumentar a complexidade do hipocampo artificial para armazenar várias lembranças e gravá-las em diversos animais, chegando até aos macacos. Planejam também poder usar essa tecnologia sem fio, substituindo os fios por rádios minúsculos, de maneira que as

lembranças possam ser baixadas por meio remoto, sem a necessidade de eletrodos incômodos implantados no cérebro.

Dado que o hipocampo está ligado ao processamento da memória em humanos, os cientistas anteveem um grande potencial de aplicações no tratamento de derrames, demência, mal de Alzheimer, e muitos outros casos decorrentes de danos ou deterioração dessa região cerebral.

Certamente, será preciso transpor muitas barreiras. Apesar de tudo o que aprendemos sobre o hipocampo desde o caso de HM, ainda há uma caixa preta com conteúdo desconhecido. Sendo assim, não é possível construir uma lembrança a partir do zero, mas se a memória da realização de uma tarefa foi processada, já é possível gravá-la e reproduzi-la.

RUMOS FUTUROS

Trabalhar com o hipocampo de primatas e humanos será mais difícil, pois é muito maior e mais complexo. O primeiro passo é criar um mapa neural detalhado da área. Isso significa colocar eletrodos em diferentes partes do hipocampo para registrar os sinais trocados constantemente entre diferentes regiões e estabelecer o fluxo de informações que corre o tempo todo pelo hipocampo. Essa região do cérebro tem quatro divisões básicas, de CA1 a CA4, e os cientistas precisam gravar os sinais trocados entre eles.

No segundo passo, alguém realiza certas tarefas, e os cientistas gravam os impulsos que fluem através das várias regiões do hipocampo, registrando a lembrança. Por exemplo: a lembrança de determinada tarefa, como saltar por dentro de uma argola, cria uma atividade elétrica no hipocampo que pode ser gravada e minuciosamente analisada. Depois é possível criar um dicionário de correspondências entre a lembrança e o fluxo das informações pelo hipocampo.

Por fim, o terceiro passo consiste em fazer um registro dessa lembrança e inserir os sinais elétricos no hipocampo de outra pessoa, via eletrodos, para ver se a lembrança foi exportada. Dessa maneira, é possível aprender a saltar por uma argola sem nunca ter feito isso antes. Se der certo, os cientistas poderão criar gradualmente uma biblioteca contendo registros de certas lembranças.

Pode levar décadas até conseguirem chegar à memória humana, mas já é possível imaginar como isso acontecerá. No futuro, talvez pessoas sejam contratadas para criar certas lembranças, como uma viagem de luxo ou uma batalha fictícia. Serão colocados nanoeletrodos em vários lugares do cérebro para registrar a lembrança. Esses eletrodos devem ser extremamente pequenos a fim de não interferirem na formação da memória.

As informações desses eletrodos serão enviadas sem fio para um computador,

e então gravadas. Assim, se alguém quiser adquirir a experiência contida nessa lembrança terá eletrodos similares colocados no seu hipocampo e a memória será inserida em seu cérebro.

É claro que essa ideia tem complicações. Se tentarmos inserir a lembrança de uma atividade física, como artes marciais, teremos o problema da “memória muscular”. Por exemplo: quando andamos, não pensamos conscientemente em pôr uma perna na frente da outra. Andar se tornou natural para nós porque fazemos isso desde muito cedo. Assim, os sinais que controlam nossas pernas já não se originam inteiramente no hipocampo, mas também no córtex motor, no cerebelo, e nos gânglios basais. No futuro, se quisermos inserir memórias ligadas a esporte, os cientistas terão que decifrar como as lembranças são também parcialmente armazenadas em outras áreas cerebrais.

VISÃO E MEMÓRIA HUMANA

A formação de memória é muito complexa, mas a abordagem que temos discutido pega um atalho, observando os sinais que se movem pelo hipocampo, onde os impulsos sensoriais já foram processados. Em *Matrix*, porém, um eletrodo é colocado na nuca para baixar memórias diretamente no cérebro. Isso pressupõe que é possível decodificar os impulsos diretos, não processados, vindos dos olhos, ouvidos, pele etc., que sobem pela medula espinhal e pelo tronco cerebral até o tálamo. É algo muito mais elaborado e difícil do que analisar as mensagens processadas circulando no hipocampo.

Para dar uma noção do volume total de informações não processadas que vêm pela medula espinhal até o tálamo, vamos considerar apenas um aspecto: a visão, já que muitas lembranças são codificadas assim. Há cerca de 130 milhões de células na retina, chamadas cones e bastonetes, que processam e registram 100 milhões de bits de informação do ambiente o tempo todo.

Essa grande quantidade de dados é colhida e enviada para o nervo ótico, que transporta 9 milhões de bits de informação por segundo para o tálamo. De lá, as informações chegam ao lobo occipital, na parte posterior do cérebro. O córtex visual, por sua vez, começa o árduo processo de analisar essa montanha de informações. O córtex visual consiste em áreas na parte de trás do cérebro, cada uma delas destinada a uma tarefa específica. São rotuladas de V1 a V8.

A área V1 é singularmente igual a uma tela; ela de fato cria um padrão na parte de trás do cérebro muito similar, em formato e configuração, à imagem original. Essa imagem tem uma notável semelhança com a original, exceto que o centro do olho, a fôvea, ocupa uma área muito maior em V1 (dado que a fôvea tem a mais alta concentração de neurônios). Assim, a imagem formada na V1

não é uma réplica perfeita, mas distorcida, com a região central da imagem ocupando a maior parte do espaço.

Além da V1, as outras áreas do lobo occipital processam diferentes aspectos da imagem, inclusive:

- Visão estéreo. Esses neurônios comparam as imagens que veem de cada olho. Isso ocorre na área V2.
- Distância. Esses neurônios calculam a distância do objeto, usando sombras e outras informações dos dois olhos. Isso ocorre na área V3.
- Cores são processadas na área V4.
- Movimento. Diferentes circuitos captam diferentes classes de movimento, inclusive linha reta, espiral e movimentos de expansão. Isso ocorre na área V5.

Foram identificados mais de 30 circuitos neurais diferentes envolvidos na visão, mas provavelmente há muitos mais.

Do lobo occipital as informações são enviadas para o córtex pré-frontal, onde finalmente “vemos” a imagem e formamos a memória de curto prazo. Em seguida as informações são enviadas para o hipocampo, onde são processadas e armazenadas por até 24 horas. Então a memória é picotada e espalhada pelos vários córtices.

A questão aqui é que a visão parece acontecer sem esforço, mas exige bilhões de neurônios disparando em sequência, transmitindo milhões de bits de informação por segundo. É importante lembrar que recebemos sinais dos cinco órgãos dos sentidos, mais as emoções associadas a cada imagem. Toda essa informação é processada pelo hipocampo para criar a simples lembrança de uma imagem. Atualmente, não existe máquina que possa se comparar à sofisticação desse processo, portanto, replicá-lo representa um enorme desafio para os cientistas que querem criar um hipocampo artificial para o cérebro humano.

LEMBRANÇAS DO FUTURO

Se codificar a memória de apenas um dos sentidos já é um processo complexo, como se desenvolveu a capacidade de armazenar grandes quantidades de informação na memória de longo prazo? Na maioria das vezes, o instinto guia o comportamento dos animais, que parecem não ter muita memória de longo

prazo. Mas, como diz o dr. James McGaugh, da Universidade da Califórnia em Irvine, “o objetivo da memória é prever o futuro”, o que levanta uma possibilidade interessante. Talvez a memória de longo prazo tenha evoluído *porque* era útil para simular o futuro. Em outras palavras, o fato de podermos recordar o passado distante deve-se a exigências e vantagens de simular o futuro.

De fato, as varreduras cerebrais feitas pelos cientistas da Universidade de Washington em St. Louis indicam que as áreas usadas para recordar são as mesmas envolvidas em simular o futuro. Em particular, a ligação entre o córtex pré-frontal dorsolateral e o hipocampo se ativa quando alguém está ocupado em planejar o futuro e lembrar o passado. Em certo sentido, o cérebro está tentando “se lembrar do futuro”, baseando-se em lembranças do passado para determinar como algo vai evoluir no futuro. Isso pode explicar também o curioso fato de que pessoas que sofrem de amnésia – como HM – são geralmente incapazes de imaginar o que farão no futuro, ou até mesmo no dia seguinte.

“Pode-se pensar nisso como uma viagem mental no tempo – a capacidade de formar pensamentos sobre si mesmo e projetá-los no passado ou no futuro”, diz a dra. Kathleen McDermott, da Universidade de Washington. Ela observa também que seu estudo prova uma “resposta possível para a questão da utilidade evolutiva da memória. Talvez a razão pela qual recordamos o passado com detalhes nítidos seja que esse conjunto de processos é importante para podermos nos visualizar em situações futuras. Essa capacidade de imaginar o futuro tem um significado adaptativo claro e decisivo”. Para um animal, o passado é um grande desperdício de recursos preciosos, pois rende pouca vantagem evolutiva. A simulação do futuro utilizando as lições do passado é uma razão essencial para os humanos terem se tornado inteligentes.

UM CÓRTEX ARTIFICIAL

Em 2012, os mesmos cientistas do Wake Forest Baptist Medical Center e da Universidade do Sul da Califórnia que criaram um hipocampo artificial em ratos anunciaram um experimento ainda mais abrangente. Em vez de gravar a memória do rato no hipocampo, eles duplicaram o processo, muito mais sofisticado, do córtex de um primata.

Pegaram cinco macacos rhesus e inseriram pequenos eletrodos em duas camadas do córtex chamadas L2/3 e L5. Registraram os sinais neurais que ocorriam entre essas duas camadas enquanto os macacos aprendiam uma tarefa. (A tarefa consistia em apresentar um conjunto de imagens para os macacos, que recebiam uma recompensa quando conseguiam identificar as mesmas imagens num escopo muito maior.) Após praticarem, os macacos conseguiram

desempenhar a tarefa com 75% de acertos. Mas se os cientistas introduzissem novamente os sinais no córtex enquanto eles cumpriam a tarefa, havia um aumento de 10% no desempenho dos macacos. Quando certas substâncias químicas eram dadas aos animais, o desempenho caía 20%, e se os registros fossem novamente inseridos no córtex, o desempenho era acima do nível normal. Apesar de ser uma amostra muito pequena, e de ter havido apenas uma melhora modesta no desempenho, o estudo sugere que os registros feitos pelos cientistas captaram com precisão o processo de tomada de decisão do córtex.

Como esse estudo foi realizado com primatas, e não com ratos, e envolveu o córtex, e não o hipocampo, pode ter grandes implicações quando começarem os testes com humanos. O dr. Sam A. Deadwyler, da Wake Forest, diz: “A ideia é que o equipamento possa gerar um padrão de saída que não passe pela área lesionada, fornecendo uma conexão alternativa” no cérebro. Esse experimento tem aplicação possível para pacientes com lesão no neocórtex. Como uma muleta, esse equipamento de apoio poderá realizar a operação de pensamento da área lesionada.

UM CEREBELO ARTIFICIAL

Vale notar também que o hipocampo e o neocórtex artificiais são apenas os primeiros passos. Em algum momento do futuro, outras partes do cérebro terão contrapartes artificiais. Por exemplo: cientistas da Universidade de Tel Aviv, em Israel, já criaram um cerebelo artificial para um rato. O cerebelo é uma parte essencial do cérebro reptiliano que controla o equilíbrio e outras funções corporais básicas.

Geralmente, quando um jato de ar é dirigido ao focinho de um rato, ele pisca. Se for feito um som ao mesmo tempo, o rato pode ser condicionado a piscar, bastando ouvir o som. A meta dos cientistas israelenses era criar um cerebelo artificial que duplicasse esse feito.

Primeiro, os cientistas registraram os sinais entrando no tronco encefálico quando o rato recebia o jato de ar e ouvia o som. Depois o sinal era processado e reenviado para o tronco cerebral em outro local. Como se esperava, os ratos piscavam ao receber o sinal. Esta foi a primeira vez que um cerebelo artificial funcionou corretamente, e foi a primeira vez que mensagens foram recebidas de uma parte do cérebro, processadas, e depois inseridas numa outra parte.

Ao comentar esse trabalho, Francesco Sepulveda, da Universidade de Essex, diz: “Isso demonstra quão longe chegamos na criação de circuitos que podem um dia substituir áreas cerebrais lesadas, e até aumentar o poder de um cérebro saudável.”

Ele vê também um grande potencial para cérebros artificiais no futuro: “Provavelmente levaremos várias décadas para chegar lá, mas meu palpite é que partes de cérebro específicas, bem organizadas, como o hipocampo ou o córtex visual, terão correlatos sintéticos antes do fim do século.”

Embora o progresso na criação de peças para a restituição do cérebro venha tendo notável rapidez em vista da complexidade do processo, é uma corrida contra o tempo considerando-se a maior ameaça ao sistema de saúde pública, que é o declínio mental de pessoas com mal de Alzheimer.

MAL DE ALZHEIMER – O DESTRUIDOR DA MEMÓRIA

Alguns dizem que Alzheimer é o mal do século. Atualmente 5,3 milhões de norte-americanos sofrem desse mal, e prevê-se que esse número quadruplique até 2050. Cinco por cento das pessoas com idade entre 65 e 74 anos têm a doença de Alzheimer, e mais de 50% dos idosos com mais de 85 anos a têm, mesmo sem apresentar fatores óbvios de risco. (Em 1900, a expectativa de vida nos Estados Unidos era de 49 anos, por isso o Alzheimer não era um problema significativo. Mas hoje a população com mais de 80 anos é um dos grupos demográficos que mais crescem no país.)

No início da doença, o hipocampo, a parte do cérebro onde a memória é processada, começa a se deteriorar. De fato, as varreduras mostram nitidamente que o hipocampo de pacientes com Alzheimer encolhe, e os filamentos entre o córtex pré-frontal e o hipocampo também ficam mais finos, deixando o cérebro incapaz de processar adequadamente a memória de curto prazo. A memória de longo prazo, já armazenada nos córtices do cérebro, permanece relativamente intacta, pelo menos no começo. Isso cria uma situação em que a pessoa pode não se lembrar do que fez cinco minutos atrás, mas se recorda claramente das décadas passadas.

A doença vai progredindo até que as lembranças básicas de longo prazo também são destruídas. O paciente não consegue reconhecer filhos nem cônjuge, não se lembra de quem ele mesmo é, e pode até entrar num estado vegetativo semelhante ao coma.

Infelizmente, os mecanismos básicos do mal de Alzheimer só começaram a ser entendidos muito recentemente. Uma grande descoberta ocorreu em 2012, quando se soube que o Alzheimer começa com a produção de proteínas amiloides tau, que estimulam a formação da proteína beta-amiloide, uma substância gosmenta, grudenta, que entope o cérebro. Até então, não se sabia ao certo se o mal de Alzheimer era causado por essas placas ou se essas placas eram subprodutos de algum outro distúrbio.

O que torna difícil combater as placas de amiloides com medicamentos é que elas devem ser formadas por “príons”, que são moléculas de proteínas deformadas. Não são bactérias nem vírus, mas mesmo assim se reproduzem. Uma molécula de proteína tem a aparência de um emaranhado de fitas de átomos entrelaçadas. Esse amontoado de átomos precisa se dobrar corretamente para a proteína assumir a forma e a função adequadas. Os príons são proteínas deformadas, que se dobraram incorretamente. O pior é que, quando esbarram em proteínas saudáveis, fazem com que elas se dobrem incorretamente também. Portanto, um príon pode causar uma cascata de proteínas deformadas, criando uma reação em cadeia que contamina bilhões de outras.

Atualmente não se conhece um modo de interromper a implacável progressão do mal de Alzheimer. Entretanto, agora que os mecanismos básicos do Alzheimer estão sendo revelados, um método promissor é criar anticorpos ou uma vacina que aja especificamente nessas moléculas de proteínas deformadas. Outro método pode ser criar um hipocampo artificial para esses indivíduos, a fim de restaurar a memória de curto prazo.

Outra abordagem seria usar a genética para aumentar diretamente a capacidade cerebral de criar memória. Talvez existam genes que possam melhorar nossa memória. O futuro da pesquisa sobre memória pode estar no “rato inteligente”.

O RATO INTELIGENTE

Em 1999, o dr. Joseph Tsien e outros cientistas de Princeton, do MIT, e da Universidade de Washington descobriram que um único gene extra acrescentado a um rato aumentava drasticamente sua memória e sua capacidade. Os “ratos inteligentes” percorriam labirintos com maior rapidez, lembravam-se melhor dos eventos, e superavam outros ratos em diversos testes. Foram apelidados de “ratos Doogie” por causa do esperto personagem da série de TV *Tal pai, tal filho*.

O dr. Tsien começou analisando o gene NR2B, que age como um interruptor no controle da capacidade do cérebro para associar um evento a outro. (Os cientistas sabem disso porque, quando o gene é silenciado ou desativado, o rato perde essa habilidade.) Todo aprendizado depende do NR2B, porque ele controla a comunicação entre as células da memória e do hipocampo. Primeiro, o dr. Tsien criou uma linhagem de ratos sem NR2B, que apresentaram deficiências de memória e dificuldade de aprendizagem. Depois ele criou uma linhagem de ratos com quantidade de genes NR2B acima do normal, que apresentaram capacidade mental superior. Colocados numa bacia rasa com água e forçados a nadar, os ratos normais nadavam aleatoriamente. Não se lembravam que havia

uma plataforma escondida na água, o que lhes foi mostrado dias antes. Os ratos espertos, porém, iam direto para a plataforma logo na primeira vez.

Desde então, pesquisadores têm confirmado esses resultados em outros laboratórios e criado linhagens de ratos ainda mais inteligentes. Em 2009, o dr. Tsien publicou um artigo anunciando mais uma linhagem de ratos inteligentes, apelidados de “Hobbie-J” (nome de um personagem de uma história em quadrinhos chinesa). Hobbie-J era capaz de lembrar fatos novos (como a localização de brinquedos) por um tempo três vezes maior do que a linhagem de ratos geneticamente modificada considerada a mais inteligente até então. “Isso reforça a noção de que NR2B é um comutador universal para formação de memória”, observou o dr. Tsien. “É como fazer de um Michael Jordan um super-Michael Jordan”, disse o estudante Deheng Wang.

No entanto, há limites, até para essa nova linhagem de ratos. Quando colocados para escolher entre virar à direita ou à esquerda para obter um pedaço de chocolate, Hobbie-J se lembrou do caminho por muito mais tempo que os ratos normais, mas depois de cinco minutos ele também esqueceu. “Nunca conseguiremos torná-lo um matemático. Afinal, ele é um rato”, disse Tsien.

É importante notar também que algumas linhagens de ratos inteligentes eram muito tímidas em comparação com ratos normais. Há suspeitas de que, se a memória é boa demais, os fracassos e mágoas também são lembrados, e isso talvez cause hesitação. Portanto, há um possível ponto negativo em se lembrar demais.

Agora, os cientistas esperam ampliar a pesquisa com testes em cães, pois compartilhamos muitos genes com eles, e talvez depois cheguem aos humanos.

MOSCAS INTELIGENTES E RATOS BOBOS

O gene NR2B não é o único sendo estudado pelos cientistas por ter impacto na memória. Em uma outra série de experimentos, os cientistas conseguiram criar uma linhagem de moscas-das-frutas com “memória fotográfica”, e uma linhagem de ratos amnésicos. Esses experimentos poderão desvendar muitos mistérios da nossa memória de longo prazo, como por que virar a noite estudando para uma prova não é a melhor maneira de aprender, ou por que nos lembramos de acontecimentos que tiveram emoções fortes. Os cientistas descobriram que há dois genes importantes, o CREB ativador (que estimula a formação de novas conexões entre neurônios), e o CREB repressor (que suprime a formação de novas memórias).

Os pesquisadores Jerry Yin e Timothy Tully, da Cold Spring Harbor, fizeram experimentos interessantes com moscas-das-frutas. Normalmente, são

necessárias dez tentativas para que aprendam uma determinada tarefa (por exemplo, detectar um odor, evitar levar um choque). As moscas com um gene CREB repressor a mais não conseguiam formar nenhuma memória duradoura, e a surpresa veio quando testaram moscas-das-frutas com um gene CREB ativador a mais. Elas aprenderam a tarefa numa única sessão. “Isso indica que essas moscas têm memória fotográfica”, disse Tully. E o cientista contou que são como aqueles alunos “que leem o capítulo de um livro uma vez, o retêm na mente, e dizem que a resposta está no parágrafo três da página duzentos e setenta e quatro”.

Tal resultado não se restringe apenas às moscas-das-frutas. O dr. Alcino Silva, também da Cold Spring Harbor, fez experimentos com ratos e descobriu que aqueles com deficiência de CREB ativador eram praticamente incapazes de formar memórias de longo prazo. Eram ratos amnésicos. Mas mesmo esses ratos esquecidos conseguiam aprender um pouco se as lições fossem curtas, e com intervalos para descanso. Os cientistas teorizam que temos no cérebro uma quantidade determinada de CREB ativador que pode limitar o volume do que podemos aprender num período específico. Se tentamos virar a noite estudando antes de uma prova, exaurimos a quantidade de CREB ativador, e não conseguimos aprender mais nada – pelo menos até fazer um intervalo para reabastecer o CREB ativador.

“Podemos dar uma razão biológica que explica por que virar a noite estudando não funciona”, diz o dr. Tully. A melhor maneira de se preparar para uma prova final é rever a matéria mentalmente, periodicamente, durante o dia, até que se torne parte da memória de longo prazo.

Isso pode explicar também por que lembranças carregadas de emoção são tão nítidas e podem durar décadas. O gene CREB repressor é como um filtro, descarta as informações inúteis. Mas se uma lembrança está associada a uma forte emoção, é capaz de remover o gene CREB repressor ou aumentar os níveis do CREB ativador.

Podemos esperar outras descobertas para o entendimento da base genética da memória. As imensas habilidades do cérebro não são formadas com apenas um gene, mas, provavelmente, com uma sofisticada combinação deles. Esses genes, por sua vez, têm contrapartes no genoma humano. Portanto, há uma real possibilidade de aumentarmos geneticamente a capacidade de nossa mente e memória.

Contudo, não espere conseguir um cérebro aditivado tão cedo. Ainda há muitos obstáculos. Primeiro, não está claro se esses resultados se aplicam a humanos. Muitas vezes, as terapias promissoras em ratos não se aplicam a outras espécies. Segundo, mesmo que esses resultados possam ser aplicados a humanos, não sabemos qual será o impacto. Por exemplo: esses genes podem melhorar a memória, mas não afetam a inteligência em geral. Terceiro, a terapia de genes

(isto é, consertar genes ruins) é mais difícil do que se pensava. Poucas doenças genéticas podem ser curadas com esse método. Ainda que os cientistas usem vírus inócuos para infectar células com o gene “bom”, o corpo vai mandar anticorpos para atacar o intruso, frequentemente tornando a terapia inútil. É possível que a inserção de um gene para melhorar a memória tenha o mesmo destino. Além disso, o campo da terapia de genes sofreu um golpe duro alguns anos atrás, quando um paciente morreu na Universidade da Pensilvânia durante o procedimento. O trabalho de modificação genética em humanos, portanto, enfrenta muitas questões éticas e legais.

Assim, os experimentos com humanos vão progredir muito mais lentamente do que com animais. Mas é possível imaginar que um dia o procedimento será aperfeiçoado e se tornará uma realidade. E então, para alterar nossos genes não será preciso mais que uma injeção no braço. Um vírus inócuo vai entrar no sangue e infectar as células normais com o próprio gene. Depois que o “gene inteligente” se incorpora às células, ele se ativa e libera proteínas que afetam o hipocampo e a formação da memória, aumentando a capacidade cognitiva e de memória.

Se a inserção de genes for muito difícil, outra possibilidade é inserir diretamente no corpo as proteínas adequadas, evitando o uso da terapia de genes. Em vez de tomar uma injeção, tomaríamos uma pílula.

A PÍLULA DA INTELIGÊNCIA

Um dos objetivos dessa pesquisa é criar uma “pílula da inteligência”, que potencialize a concentração, melhore a memória, e talvez aumente a inteligência. Laboratórios farmacêuticos já testaram vários medicamentos, como MEM 1003 e MEM 1414, que parecem melhorar as funções mentais.

Em estudos com animais, os cientistas descobriram que a memória de longo prazo é possível por meio da interação de enzimas e genes. O aprendizado ocorre quando certos caminhos neurais são reforçados pela ativação de genes específicos, como o CREB, que libera a proteína correspondente. De modo geral, quanto mais as proteínas CREB circulam no cérebro, mais rapidamente a memória de longo prazo é formada. Isso foi verificado em estudos com moluscos marinhos, moscas-das-frutas e ratos. A propriedade principal do MEM 1414 é acelerar a produção de proteínas CREB. Em testes de laboratório, animais idosos que receberam MEM 1414 formaram memórias de longo prazo com uma rapidez significativamente maior que o grupo de controle.

Os cientistas estão começando também a isolar a bioquímica exata requerida na formação de memória de longo prazo, tanto no nível genético como no

molecular. Quando o processo de formação de memória estiver bem compreendido, eles poderão desenvolver terapias para acelerar e fortalecer esse processo fundamental. Não só pacientes idosos ou com Alzheimer, mas também pessoas em geral poderão usufruir dessa “melhora cerebral”.

MEMÓRIAS PODEM SER APAGADAS?

O mal de Alzheimer pode destruir lembranças de forma indiscriminada, mas que tal apagá-las seletivamente? A amnésia é um dos temas preferidos em Hollywood. Em *A identidade Bourne*, Jason Bourne (interpretado por Matt Damon), um experiente agente da CIA, é encontrado flutuando na água, abandonado à morte. Quando volta a si, apresenta uma grave perda de memória. Ele está sendo perseguido por assassinos que querem matá-lo, mas sequer sabe quem ele mesmo é, o que aconteceu, nem por que querem matá-lo. A única pista para a memória é sua estranha habilidade de lutar instintivamente como um agente secreto.

Há documentação suficiente para mostrar que a amnésia pode ocorrer devido a acidentes, em decorrência de um trauma como um forte golpe na cabeça. Mas a memória pode ser apagada seletivamente? No filme *Brilho eterno de uma mente sem lembranças*, estrelado por Jim Carrey, um casal se conhece por acaso num trem e sente imediatamente uma atração mútua. Depois ficam chocados ao descobrir que tinham sido amantes anos atrás, mas não se lembravam disso. Descubrem que tinham pago a uma empresa para remover as lembranças um do outro após uma briga feia. Aparentemente, o destino estava lhes dando uma segunda chance para o amor.

A amnésia seletiva foi levada a um nível surpreendente em *MIB – Homens de preto*, em que Will Smith interpreta um agente de uma obscura organização que usa um “neutralizador” para apagar lembranças inconvenientes de OVNI e encontros com alienígenas. Tem até um marcador para indicar o momento inicial em que a memória deve ser apagada.

Todos esses elementos rendem histórias de suspense e recordes de bilheteria, mas alguma dessas coisas será realmente possível, mesmo no futuro?

Sabemos que a amnésia é possível, e que existem dois tipos básicos, dependendo de qual memória é afetada, se a de curto prazo ou de longo prazo. A “amnésia retrógrada” ocorre quando há um trauma ou lesão no cérebro e as lembranças preexistentes são perdidas, geralmente no evento que causou a amnésia. É uma amnésia semelhante à de Jason Bourne, que perdeu todas as lembranças anteriores ao ser abandonado no mar. Aqui o hipocampo ainda está intacto, portanto novas memórias podem ser formadas, mesmo com a memória

de longo prazo danificada. “Amnésia anterógrada” ocorre quando a memória de curto prazo é danificada, e a pessoa tem dificuldade de formar novas lembranças após o evento causador da amnésia. Em geral, a amnésia dura de alguns minutos a algumas horas, devido à lesão do hipocampo. A amnésia anterógrada foi apresentada magistralmente no filme *Amnésia*, em que um homem está determinado a vingar a morte da esposa. O problema é que sua memória só dura uns quinze minutos, por isso ele passa o tempo todo escrevendo em pedaços de papel, em fotos, e até faz tatuagens no próprio braço para se lembrar das pistas que descobriu sobre o assassinato. Lendo essas mensagens escritas para si mesmo, ele consegue acumular provas cruciais do que tinha esquecido.

O caso aqui é a perda da memória na ocasião do trauma ou doença, o que torna a amnésia seletiva de Hollywood altamente improvável. Filmes como *MIB – Homens de preto* se baseiam no argumento de que as lembranças são armazenadas sequencialmente, como num disco rígido, como se bastasse clicar na tecla “apagar” escolhendo o ponto inicial. Sabemos, porém, que as memórias são realmente fragmentadas, e as partes separadas são guardadas em diferentes locais do cérebro.

REMÉDIO PARA ESQUECER

Enquanto isso, os cientistas estão estudando certas drogas que podem apagar lembranças traumáticas que insistem em nos perturbar. Em 2009, cientistas holandeses liderados pelo dr. Merel Kindt anunciaram a descoberta de novos usos para um velho remédio, chamado propranolol, que poderia agir como uma droga “milagrosa” para aliviar a dor associada a lembranças traumáticas. Esse remédio não provocava a amnésia iniciada em certo ponto do tempo, mas tornava a dor mais tolerável – e, segundo o estudo, em apenas três dias.

A descoberta provocou uma onda de manchetes referindo-se a vítimas de transtorno de estresse pós-traumático (TEPT). Todos, de veteranos de guerra a vítimas de abuso sexual ou de acidentes terríveis, sentiram alívio dos sintomas. Mas o remédio também parecia desafiar as pesquisas sobre o cérebro, que mostram que a memória de longo prazo fica codificada, não eletricamente, mas no nível de moléculas de proteínas. Experimentos recentes, porém, sugerem que relembrar exige a recuperação e a remontagem da lembrança, de modo que a estrutura da proteína realmente poderia ser rearranjada nesse processo. Em outras palavras, a recordação modifica a lembrança. Talvez seja esta a razão pela qual o remédio funciona: o propranolol é conhecido por interferir na absorção de adrenalina, que é fundamental para criar as lembranças nítidas e duradouras resultantes de eventos traumáticos. “O propranolol se instala na célula

nervosa, e a bloqueia. Assim, a adrenalina, mesmo presente, não consegue agir”, diz o dr. James McGaugh, da Universidade da Califórnia em Irvine. Em outras palavras, sem adrenalina, a lembrança se esvai.

Testes controlados aplicados em indivíduos com lembranças traumáticas apresentaram resultados muito promissores. Mas o remédio encontrou uma barreira na questão ética de apagar memória. Alguns especialistas em ética não contestaram sua eficácia, mas desaprovaram a ideia de um remédio para esquecer, pois a memória tem o propósito de nos ensinar as lições da vida. Argumentam que mesmo as lembranças desagradáveis servem a um objetivo maior. O remédio foi condenado pelo presidente do Conselho de Bioética. Um relatório concluiu que “dessensibilizar nossa memória para fatos terríveis [iria] nos deixar confortáveis demais no mundo, indiferentes a sofrimentos, injustiças e crueldade. (...) É possível ficar insensível às piores tristezas da vida sem ficar insensível também às suas maiores alegrias?”.

O dr. David Magus, do Centro de Ética Biomédica da Universidade de Stanford, diz: “Por mais que nossos rompimentos, nossos relacionamentos sejam sofridos, aprendemos com essas experiências dolorosas. Elas nos tornam pessoas melhores.”

Outros discordam. O dr. Roger Pitman, da Universidade de Harvard, diz que, se um médico encontra uma vítima de acidente com dores insuportáveis, “deve privá-lo de morfina porque pode estar privando-o da experiência emocional completa? Quem iria defender isso? Por que a psiquiatria seria diferente? Acredito que por trás desse argumento se esconde a noção de que distúrbios mentais não são a mesma coisa que distúrbios físicos”.

O resultado desse debate poderá ter influência direta na próxima geração de medicamentos, pois o propranolol não é o único envolvido.

Em 2008, dois grupos independentes trabalhando com animais anunciaram outros medicamentos que poderiam realmente apagar memórias, e não apenas tratar o sofrimento que elas causam. O dr. Joe Tsien, do Medical College of Georgia, e seus colegas de Xangai afirmaram ter eliminado a memória em ratos usando uma proteína chamada CaMKII, e cientistas do SUNY Downstate Medical Center no Brooklyn, em Nova York, descobriram que a molécula PKMzeta também pode apagar memórias. O dr. Andre Fenson, um dos autores desse segundo estudo, disse: “Se outros trabalhos confirmarem esse parecer, algum dia veremos terapias baseadas no apagamento de memória por PKMzeta.” Além de apagar memórias, essa droga também “poderia ser útil no tratamento de depressão, ansiedade generalizada, fobias, estresse pós-traumático e dependência de drogas”, acrescentou.

Até o momento a pesquisa se limitou a animais, mas em breve começarão testes com humanos. Se os resultados com humanos forem os mesmos, a pílula do esquecimento poderá virar realidade. Não será como as pílulas de filmes de

Hollywood (que criam uma amnésia conveniente, numa época precisa, no momento oportuno), mas poderá ter grandes aplicações médicas no mundo real para pessoas perseguidas por lembranças traumáticas. Resta saber, porém, quanto seletivo será esse apagamento de memória em humanos.

O QUE PODE DAR ERRADO?

Vai chegar o dia em que poderemos gravar *todos* os sinais que passam pelo hipocampo, tálamo e pelo resto do sistema límbico, e fazer um registro confiável. Então, inserindo essa informação no cérebro, poderemos experimentar a totalidade do que outra pessoa já experimentou. A questão é: o que pode dar errado?

De fato, as implicações dessa ideia foram exploradas num filme estrelado por Natalie Wood, *Projeto Brainstorm* (1983), muito avançado para a época. No filme, cientistas criam o Hat, um capacete cheio de eletrodos que grava fielmente todas as sensações de alguém naquele momento. Quando uma pessoa passa a fita gravada para o cérebro, tem exatamente a mesma experiência sensorial. De brincadeira, um dos personagens coloca o Hat enquanto tem relações sexuais e grava a experiência. Depois a fita é regravada repetidas vezes, ampliando a experiência. E quando outra pessoa, desavisada, insere a experiência no cérebro, quase morre por causa da sobrecarga sensorial. Mais tarde uma cientista tem um infarto fatal, mas antes de morrer ela grava seus momentos finais. Quando um homem põe no cérebro a gravação da morte, ele também tem um infarto e morre.

A informação sobre essa máquina sensacional vaza e o Exército tenta assumir o controle. Isso desencadeia uma luta pelo poder entre os militares, que veem ali uma arma poderosa, e os cientistas que querem usá-la para desvendar os segredos da mente.

O *Projeto Brainstorm*, profeticamente, pôs em evidência não só a promessa dessa tecnologia, mas também os possíveis problemas. Trata-se de uma obra de ficção científica, mas alguns cientistas acreditam que, no futuro, esses mesmos problemas possam estar nas manchetes e nos tribunais de justiça.

Já vimos que houve desenvolvimentos promissores na gravação de uma única lembrança de um rato. Talvez em meados do século seja possível gravar corretamente lembranças de primatas e humanos. Mas para criar o Hat, que grava a totalidade de estímulos entrando no cérebro, é preciso registrar os dados sensoriais brutos que sobem pela medula espinhal até o tálamo. Talvez no fim do século isso aconteça.

QUESTÕES SOCIAIS E LEGAIS

Alguns aspectos desse dilema podem se manifestar em breve. Por um lado, podemos chegar ao ponto de aprender cálculo fazendo o *upload* do conhecimento. O sistema educacional sofrerá mudanças profundas. Talvez isso fará com que os professores deem mais atenção a cada aluno em áreas do conhecimento menos técnicas, que não podem ser aprendidas com um clique. A memorização necessária para se tornar médico, advogado ou cientista também pode ser reduzida drasticamente por esse método.

Em princípio, pode até nos dar lembranças de férias que nunca aconteceram, prêmios que nunca recebemos, amantes que nunca amamos, famílias que nunca tivemos. Poderia compensar deficiências, criando lembranças perfeitas de uma vida que nunca foi vivida. Os pais vão adorar, pois poderão ensinar aos filhos lições tiradas de lembranças verdadeiras. A demanda por esse tipo de dispositivo será enorme. Alguns especialistas em ética temem que essas falsas lembranças sejam tão nítidas que as pessoas prefiram reviver vidas imaginárias a ter as experiências da vida real.

Os desempregados também poderão se beneficiar aprendendo qualificações exigidas pelo mercado de trabalho, por meio de implantes de memória. Historicamente, milhões de trabalhadores ficaram para trás cada vez que foi introduzida uma nova tecnologia, muitas vezes sem qualquer amparo. É por isso que não existem mais muitos ferreiros, nem carroceiros. Eles migraram para a indústria automobilística ou outras indústrias. Mas mudar de profissão exige muito tempo e dedicação. Se o conhecimento puder ser implantado no cérebro, haverá um impacto imediato no sistema econômico mundial, pois não será preciso desperdiçar tanto capital humano. Até certo ponto, algumas habilidades poderão ser desvalorizadas se qualquer um puder baixar memórias de um determinado conhecimento, mas isso será compensado pela maior quantidade de trabalhadores qualificados.

A indústria do turismo também será impulsionada. Uma barreira para viagens internacionais é o trabalho de aprender costumes diferentes e conversar em outra língua. Os turistas poderão ter a experiência de viver numa terra estrangeira sem se atrapalhar com a moeda ou o sistema de transportes locais. Embora seja difícil que consigamos fazer o *upload* de um idioma inteiro, com milhares de palavras e expressões, pode ser possível importar informações suficientes para conduzir uma conversa agradável.

Inevitavelmente, essas gravações de lembranças estarão nas redes sociais. No futuro, será possível gravar uma lembrança e postar na internet para milhões de pessoas terem a mesma experiência. Já falamos aqui de uma mentenet, através da qual será possível enviar pensamentos. Mas se as lembranças puderem ser gravadas e criadas, será possível também enviar experiências completas. Se um

atleta ganhar uma medalha nas Olimpíadas, por que não compartilhar a agonia e o êxtase da vitória, colocando essa lembrança na rede? Se a moda pegar, bilhões de pessoas poderão compartilhar seu momento de glória.

Crianças, que estão sempre à frente em relação a videogames e redes sociais, podem ciar o hábito de gravar experiências memoráveis e publicar na internet. Assim como tirar fotos com o celular, gravar lembranças será simples e natural para elas. Para isso, é preciso que tanto o emissor como o receptor tenham nanofios quase invisíveis conectados ao hipocampo. A informação é enviada sem fio para um servidor que converte a mensagem em sinais digitais que possam ser transmitidos pela internet. Desse modo podem existir blogs, quadros de mensagens, redes sociais e salas de bate-papo onde, em vez fotos e vídeos, são publicadas lembranças e emoções.

BIBLIOTECA DE ALMAS

Talvez também queiramos ter uma árvore genealógica de memórias. Quando procuramos registros de antepassados, só encontramos retratos unidimensionais da vida deles. Ao longo da história da humanidade, as pessoas viveram, amaram e morreram sem deixar um registro substancial dessa existência. Na maioria das vezes, encontramos somente as datas de nascimento e morte de nossos ancestrais, com pouca coisa além disso. Hoje deixamos um longo rastro de documentos eletrônicos (faturas de cartão de crédito, contas, e-mails, extratos bancários etc.). Por ser hoje um meio melhor, a rede está se tornando um repositório gigantesco de todos os documentos que descrevem nossa vida, mas ainda não diz muito do que pensamos e sentimos. Talvez num futuro distante a rede se torne uma gigantesca biblioteca, contando não só os detalhes de nossa vida, mas também de nossa consciência.

No futuro, as pessoas poderão registrar suas memórias rotineiramente, para que seus descendentes possam compartilhar da mesma experiência. Na biblioteca de memórias dos nossos antepassados, poderemos ver e sentir como eles viveram, e saber como nos encaixamos nesse contexto maior.

Isso significa que qualquer pessoa poderá reviver nossa vida, muito depois da nossa morte, bastando dar um “play”. Se essa previsão estiver correta, poderemos também “trazer de volta” nossos ancestrais para bater um papo, simplesmente inserindo um disco na biblioteca e clicando num botão.

E se quisermos compartilhar das experiências de nossos personagens históricos preferidos, será possível ver de perto o que eles sentiram ao enfrentar grandes crises. Se tivermos algum ídolo e quisermos saber como ele enfrentou e venceu grandes desafios na vida, é só experienciar as memórias gravadas e adquirir

conhecimentos valiosos. Imagine compartilhar da memória de um cientista ganhador do Prêmio Nobel. Poderemos aprender como fazer grandes descobertas. Ou poderemos compartilhar das lembranças de grandes políticos e estadistas quando tomaram decisões que mudaram a história do mundo.

O dr. Miguel Nicolelis acredita que um dia tudo isso será realidade. “Cada uma dessas gravações perenes será valorizada como uma joia preciosa, uma entre bilhões de mentes também singulares que viveram, amaram, sofreram e prosperaram. Só que elas foram imortalizadas, não confinadas em túmulos frios e silenciosos, mas libertadas por meio de pensamentos brilhantes, amores intensos e tristezas.”

O LADO NEGRO DA TECNOLOGIA

Alguns cientistas ponderaram sobre as implicações éticas dessa tecnologia. Quase todas as descobertas médicas causaram preocupações éticas quando foram introduzidas. Algumas delas tiveram que ser limitadas ou proibidas por serem prejudiciais (como o calmante talidomida, que causava má-formação de bebês). Outras tiveram tanto sucesso que mudaram nosso conceito de quem somos, como os bebês de proveta. O nascimento de Louise Brown, o primeiro bebê de proveta, em 1978, criou tamanho alvoroço na mídia que até o papa deu uma declaração criticando a tecnologia. Mas hoje, seu irmão, filho, cônjuge, ou até você mesmo, pode ser um produto da fertilização *in vitro*. Como no caso de várias tecnologias, o público acabará se acostumando com a ideia de memórias gravadas e compartilhadas.

Outros especialistas em bioética têm preocupações diferentes. O que acontece se nos dão lembranças sem nossa permissão? O que acontece se essas lembranças são dolorosas ou destrutivas? E quanto aos pacientes de Alzheimer, que se beneficiariam com *uploads* de memórias, mas estão doentes demais para dar permissão?

O falecido Bernard Williams, filósofo da Universidade de Oxford, questionou se esse procedimento não perturbaria a ordem natural das coisas, que é esquecer. “O esquecimento é o processo natural mais benéfico que possuímos”, ele disse.

A possibilidade de implantar memórias tão facilmente quanto importar arquivos para o computador pode abalar as bases do sistema legal. Se um dos pilares da justiça é a testemunha ocular, o que acontece se falsas memórias são implantadas? E se a memória de um crime pode ser criada, também pode ser plantada no cérebro de um inocente. Ou, se um criminoso precisar de um *álibi*, pode inserir secretamente a memória no cérebro de outra pessoa, convencendo-a de que estavam juntos quando o crime foi cometido. Além disso, não só os

depoimentos verbais, mas também os documentos legais serão suspeitos, porque quando assinamos declarações sob juramento dependemos da memória para esclarecer o que é verdadeiro ou falso.

Será preciso introduzir formas de proteção, criar leis que definam claramente os limites de acesso ou recusa de memórias. Assim como há leis que restringem o direito da polícia ou de terceiros a entrar em nossa casa, haverá leis restringindo o direito de acessar nossas memórias sem permissão. Será preciso também criar um meio de marcar essas memórias para percebermos que são fictícias. Assim poderemos ter o prazer de recordar uma agradável viagem de férias, mas sabendo que nunca existiu.

Acessar, armazenar e transferir memórias nos permitem registrar o passado e adquirir novas habilidades. Mas isso não altera nossa capacidade inata de digerir e processar esse volume de informações. Para tanto, precisamos expandir nossa inteligência. O progresso nessa direção encontra dificuldade por não existir uma definição universal de inteligência. No entanto, um exemplo de genialidade e inteligência que ninguém pode contestar é Albert Einstein. É impressionante que sessenta anos após sua morte seu cérebro ainda forneça pistas sobre a natureza da inteligência.

Alguns cientistas acreditam que com uma combinação de terapias eletromagnéticas, genéticas e medicamentosas, é possível elevar nossa inteligência ao nível de gênio. Eles citam o fato de que há registros de lesões cerebrais que podem fazer com que uma pessoa normal adquira subitamente as qualidades de um “savant”, que tem capacidade mental e artística espetacular, muito além do normal. Isso ocorre hoje em consequência de acidentes, mas o que pode acontecer se a ciência conseguir desvendar o segredo desse processo?

O cérebro é maior que o céu
Pois se postos lado a lado
O primeiro o outro contém
E ainda cabe você também
– EMILY DICKINSON

O talentoso acerta um alvo que ninguém consegue acertar. O gênio acerta
um alvo que ninguém consegue ver.
– ARTHUR SCHOPENHAUER

6 O CÉREBRO DE EINSTEIN E A EXPANSÃO DE NOSSA INTELIGÊNCIA

O cérebro de Einstein sumiu.

Ou pelo menos ficou sumido durante cinquenta anos, até que os herdeiros do médico, que desapareceram com o cérebro dele pouco depois de sua morte em 1955, devolveram-no ao National Museum of Health and Medicine em 2010. A análise desse cérebro pode ajudar a esclarecer algumas questões: O que é a genialidade? Como se pode medir inteligência e qual a relação dela com sucesso na vida? Há também questões filosóficas: A genialidade é função dos genes, ou é mais uma questão de esforço e realização pessoal? E, por fim, o cérebro de Einstein pode ajudar a responder a pergunta fundamental: Podemos expandir nossa própria inteligência?

A palavra “Einstein” já não é mais um nome próprio referido a uma pessoa específica. Hoje significa simplesmente “gênio”. A imagem que o nome evoca (calças frouxas, cabelos brancos desalinhados, aparência desarrumada) é igualmente icônica e instantaneamente reconhecível.

O legado de Einstein é imenso. Em 2011, quando alguns físicos sugeriram que ele estava errado, que a barreira da luz podia ser quebrada por partículas, a controvérsia no mundo da física foi tão grande que chegou à imprensa popular. A própria ideia de que a relatividade – a base da física moderna – pudesse estar errada causou indignação em físicos do mundo inteiro. Como era de se esperar, quando os resultados foram reconsiderados, viu-se que Einstein tinha razão. É sempre perigoso questionar Einstein.

Uma forma de abordar a pergunta “O que é genialidade?” é analisar o cérebro de Einstein. Aparentemente no calor do momento, o dr. Thomas Harvey, médico do Princeton Hospital que estava fazendo a autópsia de Einstein, decidiu preservar secretamente o cérebro dele, sem o conhecimento nem autorização da família.

Talvez ele tenha preservado o cérebro de Einstein com a vaga noção de que algum dia pudesse revelar o segredo da genialidade. Talvez tenha pensado, como muitos outros, que houvesse uma parte peculiar daquele cérebro que abrigasse sua enorme inteligência. Brian Burrell, no livro *Postcards from the Brain Museum*, especula que talvez o dr. Harvey tenha “se deixado levar pelo momento e ficado fora de si diante daquela grandiosidade. O que ele logo descobriu foi que tinha dado um passo maior que a perna”.

O que aconteceu com o cérebro de Einstein depois disso parece mais uma comédia do que uma história científica. Durante anos, o dr. Harvey prometeu publicar o que descobrisse sobre o cérebro de Einstein. Mas ele não era especialista em cérebro, e passou a dar desculpas. O cérebro passou décadas guardado em dois potes grandes de conserva cheios de formol, dentro de uma embalagem de sidra escondida embaixo de uma caixa térmica. Ele chamou um

técnico para cortar o cérebro em 240 partes e, em raras ocasiões, escrevia para alguns cientistas que queriam estudar o cérebro. Certa vez, alguns pedaços do cérebro foram enviados dentro de um vidro de maionese para um cientista em Berkeley.

Quarenta anos depois, Harvey atravessou o país dirigindo um Buick Skylark, levando o cérebro de Einstein num tupperware, com a intenção de devolvê-lo à neta de Einstein, Evelyn. Ela se negou a aceitar. Após a morte do dr. Harvey, em 2007, a coleção de slides e de pedaços do cérebro foi doada para estudos científicos, por decisão dos herdeiros do cientista. A história do cérebro de Einstein é tão inusitada que rendeu um documentário para a televisão.

Vale notar que o cérebro de Einstein não foi o único preservado para a posteridade. O cérebro de um dos maiores gênios da matemática, Carl Friedrich Gauss, conhecido como o “príncipe dos matemáticos”, também foi preservado por um médico, um século antes. Na época, a anatomia do cérebro era tão inexplorada que não se sabia nada além de que o cérebro apresentava dobras e circunvoluções inusitadamente grandes.

Seria de se esperar que o cérebro de Einstein fosse muito mais avançado que o dos humanos normais, que fosse enorme, talvez com áreas anormalmente grandes. Na verdade, viu-se que era o oposto (é um pouco menor, e não maior que o normal). No todo, o cérebro de Einstein é muito comum. Um neurologista que não soubesse a quem pertencia aquele cérebro não o olharia duas vezes.

As únicas diferenças encontradas no cérebro de Einstein são muito pequenas. Uma parte, chamada giro angular, é maior que o normal, com as regiões parietais inferiores dos dois hemisférios 15% mais largas que a média. Essas partes lidam com o pensamento abstrato, a manipulação de símbolos, como na escrita e na matemática, e o processamento visual e espacial. Mas continua sendo um cérebro dentro dos padrões, portanto, não determina se o gênio de Einstein está na estrutura orgânica do cérebro ou na força de sua personalidade, sua forma de ver o mundo, seu tempo. Numa biografia dele que escrevi, intitulada *O cosmo de Einstein*, ficou claro para mim que certos aspectos da vida dele foram tão importantes quanto qualquer anomalia de seu cérebro. Talvez o próprio Einstein tenha expressado isso melhor ao dizer, “não tenho talentos especiais. (...) Sou apenas apaixonadamente curioso”. Na verdade, Einstein confessou que teve dificuldade com a matemática quando era jovem, e confidenciou a um grupo de meninos: “Por mais que vocês tenham dificuldades com a matemática, as minhas foram maiores.” Então, por que Einstein foi Einstein?

Primeiro, Einstein passou a maior parte do tempo pensando via “experimentos mentais”. Era um físico teórico, e não experimental, e estava sempre fazendo simulações do futuro sofisticadas na cabeça. Em outras palavras, seu laboratório era a mente.

Segundo, era conhecido por passar dez anos ou mais num único experimento mental. Dos 16 aos 26 anos, Einstein se concentrou no problema da luz e se seria possível ultrapassar a velocidade da luz. Isso levou ao nascimento da relatividade especial, que veio a revelar o segredo das estrelas e nos deu a bomba atômica. Dos 26 aos 36 anos, ele se concentrou em uma teoria da gravidade, que acabou nos dando os buracos negros e a teoria do Big Bang do universo. E dos 36 anos até o fim da vida ele tentou encontrar uma teoria de tudo para unificar toda a física. A capacidade de passar dez ou mais anos num único problema mostra claramente a persistência com que ele simulava experimentos na cabeça.

Terceiro, sua personalidade tem muita importância. Nascido na Boêmia, era natural para ele se rebelar contra a física instituída. Nem todo físico tinha a audácia ou imaginação para contestar a teoria de Isaac Newton, dominante na época e nos duzentos anos que precederam Einstein.

Quarto, a época era propícia para o surgimento de um Einstein. Em 1905, o velho mundo físico de Newton estava se desfazendo diante de experimentos sugerindo claramente que uma nova física estava prestes a nascer, à espera de um gênio que abrisse o caminho. Por exemplo: a misteriosa substância chamada rádio brilhava no escuro e indefinidamente, como se a energia fosse criada do ar, violando a teoria de conservação da energia. Em outras palavras, Einstein era o homem certo para seu tempo. Se algum dia for possível clonar Einstein a partir das células de seu cérebro preservado, suspeito que o clone não será um novo Einstein. As circunstâncias históricas também têm que ser apropriadas para a criação de um gênio.

A questão aqui é que o gênio talvez seja uma combinação de certas capacidades mentais inatas com a determinação e o vigor para grandes conquistas. A essência da genialidade de Einstein foi, provavelmente, sua capacidade extraordinária de simular o futuro por meio de experimentos mentais, criando novos princípios físicos via imagens. Como disse o próprio Einstein, “o verdadeiro sinal de inteligência não é o conhecimento, mas a imaginação”. E, para Einstein, a imaginação significava romper as fronteiras do conhecido e penetrar nos domínios do desconhecido.

Todos nós nascemos com algumas capacidades programadas nos genes e na estrutura do cérebro. É uma loteria, questão de sorte. Mas a organização dos pensamentos, das experiências, e a simulação do futuro estão totalmente sob nosso controle. O próprio Charles Darwin escreveu: “Sempre afirmei que, à exceção dos tolos, os homens não diferem muito em intelecto, apenas em dedicação e trabalho árduo.”

AGENIALIDADE PODE SER APRENDIDA?

Isso reabre a questão: Alguém nasce gênio, ou vira gênio? Como o debate inato/adquirido soluciona o mistério da inteligência? Uma pessoa comum pode se tornar um gênio?

Como as células cerebrais têm uma notória dificuldade de crescer, pensava-se que a inteligência já estava pronta e fixa nos jovens adultos. Mas novas pesquisas deixam cada vez mais claro que, à medida que aprende, o cérebro pode sofrer modificações. As células cerebrais não vão sendo adicionadas ao córtex, mas as conexões entre os neurônios mudam cada vez que se aprende algo novo.

Por exemplo: em 2011, cientistas analisaram o cérebro dos famosos taxistas londrinos que precisam decorar as 25 mil ruas do confuso labirinto que a moderna Londres se tornou. São três ou quatro anos de preparação para esse difícil teste, e somente 50% dos alunos são aprovados.

Cientistas da University College London estudaram os cérebros desses taxistas antes do teste, e os examinaram novamente três ou quatro anos depois. Os que haviam passado no teste apresentaram um volume de massa cinzenta maior do que tinham antes, numa área chamada hipocampo posterior e anterior. Como vimos, o hipocampo é onde a memória é processada. (Curiosamente, os testes também mostraram que esses taxistas apresentaram resultados mais baixos que o normal no processamento de informação visual, portanto, talvez haja uma barganha, um preço a pagar para assimilar esse volume de informações.)

“O cérebro humano continua tendo ‘plasticidade’ mesmo na vida adulta, podendo se adaptar quando aprendemos coisas novas”, diz Eleanor Maguire, do Wellcome Trust, que patrocinou o estudo. “Isso estimula os adultos que querem novos aprendizados mais tarde na vida.”

Da mesma forma, o cérebro de ratos que aprenderam tarefas é um pouco diferente do cérebro de ratos que não aprenderam nada. Não é tanto pela mudança no número de neurônios; a natureza das conexões neurais é que foi alterada pelo processo de aprendizado. Em outras palavras, o aprendizado muda de fato a estrutura do cérebro.

Isso lembra o velho ditado “a prática leva à perfeição”. O psicólogo canadense dr. Donald Hebb fez uma descoberta importante sobre as conexões cerebrais: quanto mais exercitamos certas habilidades, mais reforçamos certos caminhos no cérebro, e portanto as tarefas se tornam mais fáceis. Ao contrário do computador, que continua sendo o mesmo ignorante de sempre, o cérebro é uma máquina que aprende, com a capacidade de reconectar seus caminhos neurais cada vez que aprende alguma coisa. Essa é uma diferença fundamental entre o computador e o cérebro.

Tal lição se aplica não só aos taxistas de Londres, mas também a músicos concertistas experientes. Segundo o psicólogo dr. K. Anders Ericsson e colegas, que estudaram mestres da elitista Academia de Música de Berlim, os melhores violinistas já tinham acumulado dez mil horas de prática aos 20 anos de idade,

estudando mais de trinta horas por semana. Em contraste, ele descobriu que alunos que eram meramente excepcionais tinham completado apenas oito mil horas ou menos, e futuros professores de música tinham praticado apenas um total de quatro mil horas. O neurologista Daniel Levitin diz: “O quadro extraído desses estudos é que são necessárias dez mil horas de prática para se atingir o grau de maestria associado a um expert em nível mundial – em qualquer coisa. (...) Em todos os estudos, seja com compositores, jogadores de basquete, romancistas, esquiadores, pianistas, enxadristas, grandes criminosos, seja o que for, sempre se chega a esse número.” Malcolm Gladwell, em seu livro *Fora de série*, chama isso de “a regra das 10 mil horas”.

COMO SE MEDE A INTELIGÊNCIA?

Mas como se mede a inteligência? Durante séculos, toda discussão sobre inteligência se baseava no boca a boca e em relatos. Mas hoje estudos com IRM mostram que a principal atividade do cérebro ao resolver problemas matemáticos envolve o caminho que conecta o córtex pré-frontal (que trata do pensamento racional) aos lobos parietais (que processam os números). Isso tem correlação com os estudos anatômicos do cérebro de Einstein, cujos lobos parietais inferiores são maiores que o normal. Assim, pode-se pensar que a capacidade matemática esteja correlacionada ao aumento dos fluxos de informações entre o córtex pré-frontal e os lobos parietais. Mas o tamanho dessa área do cérebro aumentou devido à grande atividade de trabalho e estudo, ou Einstein já nasceu assim? A resposta ainda não é clara.

O problema principal é que não há sequer uma definição amplamente aceita de inteligência, quanto mais um consenso entre os cientistas sobre a sua origem. Mas a resposta pode ser decisiva se quisermos expandir a inteligência.

OS TESTES DE Q.I. E O DR. Terman

A medida de inteligência mais usada é o teste de Q.I., introduzido pelo dr. Lewis Terman, da Universidade de Stanford, que em 1916 adaptou um teste anterior criado por Alfred Binet para o governo francês. Por muitas décadas, o teste se manteve como o critério para medir a inteligência. Terman, de fato, dedicou a vida à ideia de que a inteligência podia ser medida e herdada, e de que era o melhor prognóstico de sucesso na vida.

Cinco anos depois, Terman deu início a um estudo histórico com crianças em

idade escolar, chamado *The Genetic Studies of Genius*. Foi um projeto ambicioso, com escopo e duração sem precedentes nos anos 1920, e definiu o estilo de pesquisa nesse campo para toda uma geração. Ele acompanhou metodicamente os sucessos e fracassos desses indivíduos durante a vida deles, catalogando suas descobertas em volumosos arquivos. Os estudantes com Q.I. alto foram apelidados de “*Termites*” (cupins).

A princípio, a ideia do dr. Terman fez um enorme sucesso. Tornou-se o padrão de medida para crianças e para outros testes. Durante a Primeira Guerra Mundial, o teste foi aplicado em 1,7 milhão de soldados. Mas, com o passar dos anos, um perfil diferente começou a emergir. Algumas décadas depois, as crianças que tiveram alta pontuação no teste de Q.I. apresentaram um sucesso apenas moderado em relação às que tinham menor pontuação. Terman exibia orgulhosamente alguns de seus estudantes que vieram a receber prêmios e ter empregos bons e bem pagos. Mas ele foi ficando cada vez mais intrigado com a quantidade de alunos entre os mais brilhantes que eram considerados socialmente fracassados, tinham trabalhos subalternos e sem futuro, se envolviam em crimes ou levavam a vida à margem da sociedade. Esses resultados eram muito frustrantes para o dr. Terman, que dedicara a vida a provar que um Q.I. alto garantia sucesso.

SUCESSO NA VIDA E SATISFAÇÃO ADIADA

Uma abordagem diferente foi trazida em 1972 pelo dr. Walter Mischel, também de Stanford, que analisou outra característica em um grupo de crianças: a capacidade de adiar a satisfação. Ele introduziu o “teste do marshmallow”, isto é, a criança preferia ganhar um marshmallow imediatamente, ou dois dali a vinte minutos? Seiscentas crianças de idades entre 4 e 6 anos participaram desse experimento. Como Mischel verificou, em 1988, quando visitou os participantes, os que adiaram a satisfação eram mais competentes que os demais.

Em 1990, outro estudo mostrou uma correlação direta entre os que conseguiam adiar a satisfação e os resultados do SAT (teste utilizado no processo de seleção das universidades americanas). E uma pesquisa realizada em 2011 indicou que essa característica permanecia ao longo da vida da pessoa. Os resultados deste e de outros estudos surpreenderam os pesquisadores. As crianças capazes de ter a satisfação adiada tinham maior pontuação em quase todos os indicadores de sucesso na vida: trabalhos mais bem pagos, menos propensão ao vício em drogas, maior pontuação em testes, melhor formação e integração social etc.

No entanto, o mais interessante é que as varreduras cerebrais desses indivíduos

revelaram um padrão definido. Eles apresentaram uma diferença notável no modo de interação do córtex pré-frontal com o corpo estriado ventral, que é uma região ligada ao vício em drogas. (Isso não é de surpreender, dado que o estriado ventral contém o núcleo accumbens, conhecido como “centro do prazer”. Assim, parece haver uma luta entre a parte do cérebro que busca o prazer e a parte racional que controla a tentação, como vimos no capítulo 2.)

Essa diferença não é obra do acaso. A interação foi testada por vários grupos independentes ao longo dos anos, com resultados quase idênticos. Outros estudos também verificaram a diferença no circuito frontoestriatal, que parece governar a satisfação adiada. Aparentemente, a característica mais relacionada a sucesso na vida, que tem persistido no tempo, é a capacidade de adiar a satisfação.

Embora seja uma simplificação grosseira, essas varreduras mostram que a conexão entre os lobos pré-frontal e parietal deve ser importante para o pensamento matemático e abstrato, ao passo que a conexão entre os sistemas pré-frontal e límbico (que inclui o controle consciente das emoções e o centro de prazer) deve ser essencial para o sucesso na vida.

O dr. Richard Davidson, neurocientista da Universidade de Winsconsin-Madison, conclui que “as notas na escola, a pontuação no vestibular representam menos para o sucesso na vida do que a capacidade de cooperar, de regular as emoções, de adiar a satisfação e de se concentrar. Essas habilidades são muito mais importantes – todos os dados indicam – para o sucesso na vida, do que o Q.I. ou as notas na escola”.

NOVAS FORMAS DE MEDIR INTELIGÊNCIA

É claro que precisamos de novas formas de medir inteligência e sucesso na vida. Os testes de Q.I. não são inúteis, mas medem apenas uma forma limitada de inteligência. O dr. Michael Sweeney, autor de *Brain: The Complete Mind*, observa que “os testes não medem motivação, persistência, habilidades sociais, nem uma série de outros atributos para uma vida bem vivida”.

O problema com muitos desses testes padronizados é que pode haver algum desvio inconsciente, devido a influências culturais. Além disso, avaliam apenas uma forma particular de inteligência, que alguns psicólogos chamam de inteligência “convergente”. A inteligência convergente se concentra em apenas uma linha de pensamento, ignorando as formas de inteligência mais complexas, “divergentes”, que exigem um tipo de medição diferente. Por exemplo: durante a Segunda Guerra Mundial, a Força Aérea Americana pediu a cientistas que criassem um teste psicológico para medir a inteligência e a capacidade dos pilotos de lidar com situações difíceis, inesperadas. Uma pergunta era: Se você

for atingido e cair em território inimigo, o que você faz para sair de lá? Os resultados contradisseram o pensamento convencional.

A maioria dos psicólogos esperava que os pilotos com Q.I. alto tivessem um resultado melhor nesse teste também. Na verdade, ocorreu o inverso. Os pilotos com maior pontuação foram os que tinham nível mais elevado de pensamento divergente, os que consideravam muitas linhas de pensamento diferentes. Os pilotos que se sobressaíram eram capazes, por exemplo, de imaginar vários métodos criativos, não ortodoxos, para escapar quando capturados.

A diferença entre pensamento convergente e divergente aparece também em estudos com pacientes de cérebro dividido, mostrando claramente que as conexões em cada hemisfério são voltadas primordialmente para um ou para outro tipo de pensamento. O dr. Ulrich Kraft, de Fulda, na Alemanha, escreve: “O hemisfério esquerdo é responsável pelo pensamento convergente, e o direito, pelo divergente. O lado esquerdo examina detalhes e os processos lógicos e analiticamente, mas não se arrisca nem opera com conexões abstratas. O lado direito é mais imaginativo e intuitivo, tende a operar de modo holístico, integrando as peças de um quebra-cabeça para formar um todo.”

Neste livro, adoto a posição de que a consciência humana envolve capacidade de criar um modelo do mundo e simulá-lo no futuro, a fim de atingir metas. Os pilotos que demonstraram ter pensamento divergente foram capazes de simular vários eventos possíveis, com precisão e maior complexidade. Da mesma forma, crianças que foram capazes de adiar a satisfação no famoso teste do marshmallow parecem ser as que têm maior capacidade de simular o futuro, de ver a recompensa a longo prazo, e não ansiar por conquistas de curto prazo, do tipo ficar rico logo.

É difícil, mas não impossível, criar um teste de inteligência mais sofisticado, que meça especificamente a capacidade de simular o futuro. A pessoa é orientada a criar o maior número possível de cenários realistas do futuro para vencer um jogo, e a pontuação depende do número de simulações imaginadas e do número de relações causais envolvidas em cada simulação. Em vez de medir a capacidade de simplesmente assimilar informações, esse novo método irá medir a capacidade de manipular e adaptar essa informação para atingir um objetivo mais ambicioso. Por exemplo: pergunta-se a uma pessoa o que ela faria para fugir de uma ilha deserta cheia de animais selvagens famintos e cobras venenosas. Ele faz uma lista de todas as maneiras de sobreviver, evitar os animais perigosos e sair da ilha, criando uma elaborada árvore causal de resultados e futuros possíveis.

Assim, vemos que há um fio condutor em toda essa discussão, ou seja, a inteligência parece estar ligada à complexidade na simulação de eventos futuros, que está relacionada à nossa discussão anterior sobre consciência.

Mas em vista dos rápidos avanços nos laboratórios de todo o mundo em

pesquisas com campos eletromagnéticos, genética e terapias medicamentosas, será possível, não só medir a inteligência, mas também aumentá-la – para se tornar um outro Einstein?

POTENCIALIZANDO A INTELIGÊNCIA

Essa possibilidade foi explorada no romance *Flores para Algernon* (1958), que dez anos mais tarde deu origem ao filme *Os dois mundos de Charly*. É a triste história da vida de Charly Gordon, que tem um Q.I. 68 e um empreguinho numa padaria. Ele leva uma vida simplória, não entende que os colegas estão sempre zombando dele, e nem sabe soletrar o próprio nome.

Sua única amiga é Alice, uma professora que se compadece dele e tenta ensiná-lo a ler. Um dia, cientistas descobrem um novo procedimento que faz ratos comuns ficarem imediatamente inteligentes. Alice ouve falar nisso e apresenta Charly aos cientistas, que concordam em aplicar o procedimento a essa primeira cobaia humana. Em poucas semanas, Charly apresenta uma mudança notável. Seu vocabulário aumenta, ele devora livros na biblioteca, torna-se um sedutor, compra obras de arte moderna. Logo começa a ler sobre as teorias da relatividade e quântica, cruzando as fronteiras da física avançada. Ele e Alice inclusive se tornam amantes.

Mas depois os cientistas observam que os ratos vão perdendo lentamente as capacidades adquiridas, e morrem. Charly entende que ele também pode perder tudo aquilo e tenta furiosamente usar seu intelecto superior para encontrar a cura, mas não há outro jeito senão testemunhar o próprio declínio. Seu vocabulário vai diminuindo, ele esquece a matemática e a física, e volta gradualmente ao estado anterior. Na cena final, Alice, desolada, vê Charly brincando com crianças num parquinho.

O livro e o filme, embora pungentes e aclamados pela crítica, foram considerados pura ficção científica. A trama era emocionante e original, mas a ideia de potencializar a inteligência foi considerada absurda. As células do cérebro não se regeneram, diziam os cientistas, portanto a história do filme era obviamente impossível.

Mas não é mais.

Embora ainda seja impossível aumentar a inteligência, os rápidos avanços em sensores eletromagnéticos, genética e células-tronco podem um dia conseguir tal feito. O interesse científico está especialmente concentrado em “autistas savants”, que possuem capacidades fenomenais, sobre-humanas, que vão além da imaginação. E o mais importante é que em consequência de danos específicos no cérebro, pessoas normais podem adquirir rapidamente poderes quase

miraculosos. Alguns cientistas acreditam até que essas capacidades fantásticas podem ser induzidas com o uso de campos eletromagnéticos.

SAVANTS: SUPERGÊNIOS?

Uma bala atravessou o crânio de Z quando ele tinha 9 anos de idade. Não o matou, como os médicos temiam, mas causou uma lesão tão extensa no lado esquerdo do cérebro que o deixou surdo, mudo e com o lado direito do corpo paralisado.

Contudo, a bala teve um efeito colateral estranho. Z desenvolveu habilidades mecânicas excepcionais e uma memória prodigiosa, típica dos savants.

Z não é o único. Em 1979, um menino de 10 anos chamado Orlando Serrell ficou inconsciente ao ser atingido por uma bola de beisebol no lado esquerdo da cabeça. No início, ele se queixava de fortes dores de cabeça. Quando as dores passaram, ele se mostrou capaz de fazer cálculos matemáticos complicadíssimos, e com uma memória quase fotográfica para certos acontecimentos de sua vida. Ele sabia calcular datas de milhares de anos no futuro.

No mundo inteiro, de aproximadamente sete bilhões de pessoas, existem apenas cem casos documentados desses incríveis savants. (O número é muito maior se incluirmos pessoas com capacidades mentais extraordinárias, mas que não chegam a ser sobre-humanas. Acredita-se que cerca de 10% dos indivíduos autistas têm algumas características de savants.) Esses savants fenomenais possuem capacidades muito além do entendimento científico atual.

Há vários tipos de savants que despertaram recentemente a curiosidade dos cientistas. Cerca da metade deles tem alguma forma de autismo (a outra metade apresenta outros tipos de doença mental ou distúrbios psicológicos). Em geral, eles têm sérios problemas de interação social, o que leva a um profundo isolamento.

Existe também a “síndrome de savant adquirida”, em pessoas que parecem perfeitamente normais e sofrem um trauma extremo em algum momento da vida (por exemplo, ao bater a cabeça no fundo da piscina, ou ao ser atingido por uma bala ou uma bola de beisebol), quase sempre no lado esquerdo do cérebro. Alguns cientistas, porém, sugerem que essa distinção é enganosa, e que talvez *todas* as capacidades dos savants sejam adquiridas. Dado que os autistas savants começam a apresentar tais capacidades por volta dos 3 ou 4 anos de idade, talvez o autismo (como um golpe na cabeça) seja o desencadeador.

Os cientistas discordam sobre a origem dessas capacidades extraordinárias. Alguns acreditam que esses indivíduos simplesmente nasceram assim e, portanto,

são anomalias únicas, singulares. Suas habilidades, ainda que despertadas por uma bala, são estruturadas no cérebro desde o nascimento. Se assim for, talvez essa capacidade nunca possa ser aprendida nem transferida.

Outros afirmam que essa ideia viola a teoria da evolução, que vai sendo incrementada ao longo do tempo. Se existem savants geniais, então todos nós possuímos capacidades similares, embora em estado latente. Isso significa que um dia poderemos acionar intencionalmente esses poderes extraordinários? Alguns acreditam que sim, e há artigos publicados afirmando que algumas capacidades dos savants estão latentes em todos nós e podem ser despertadas com aplicações de campos magnéticos gerados por estimulação eletromagnética transcraniana. Ou talvez tais habilidades tenham uma base genética e, nesse caso, uma terapia com genes poderia recriá-las. E talvez também seja possível cultivar células-tronco que façam crescer neurônios no córtex pré-frontal e em outros centros importantes do cérebro. Assim, seria possível aumentar nossas capacidades mentais.

Todos esses caminhos são fonte de muita especulação e pesquisa. Podem não só levar os médicos a reverter as sequelas de doenças como o Alzheimer, mas também nos ajudar a aumentar a inteligência. As possibilidades são fascinantes.

O primeiro caso documentado de savant foi registrado em 1789, pelo dr. Benjamin Rush, após estudar um indivíduo que parecia ter uma deficiência mental. No entanto, quando lhe perguntaram quantos segundos tinha vivido um homem (que tinha 70 anos, 17 dias e 12 horas), ele levou 90 segundos para dar a resposta correta: 2.210.500.800.

O dr. Darold Treffert, médico de Wisconsin, que durante muito tempo estudou os savants, conta o caso de um cego ao qual se fez uma pergunta simples. Se você colocar um grão de milho no primeiro quadrado de um tabuleiro de xadrez, dois grãos no segundo, quatro grãos no terceiro, e for duplicando na sequência, quantos grãos terão nos 64 quadrados? Em 45 segundos, ele respondeu corretamente: 18.446.744.073.709.551.616.

Talvez o exemplo mais famoso de savant tenha sido o falecido Kim Peek, que inspirou o filme *Rain Man*, estrelado por Dustin Hoffman e Tom Cruise. Apesar de ter uma deficiência mental grave (era incapaz de viver sozinho e mal conseguia amarrar os sapatos e abotoar a camisa), Kim Peek memorizou cerca de 12 mil livros e sabia citar passagens, palavra por palavra, de qualquer página, que levava oito segundos para ser lida. (Ele conseguia decorar um livro em meia hora, mas lia de um modo muito inusitado: conseguia ler duas páginas ao mesmo tempo, cada uma com um olho.) Embora fosse incrivelmente tímido, ele passou até a gostar de exibir suas proezas matemáticas diante de espectadores curiosos, que o desafiavam com perguntas capciosas.

É claro que os cientistas precisam ter o cuidado de distinguir as verdadeiras capacidades de savants de simples fórmulas de memorização. Suas habilidades

não são apenas matemáticas, mas se estendem incrivelmente à música, às artes e à mecânica. Como os autistas savants têm grande dificuldade de expor verbalmente seus processos mentais, outro caminho é investigar indivíduos que têm síndrome de Asperger, que é uma forma mais amena de autismo. A síndrome só foi reconhecida como um estado psicológico específico em 1994, portanto existem muito poucas pesquisas sólidas nessa área. Assim como os autistas, os pacientes com Asperger têm dificuldade de interagir socialmente. Entretanto, com cuidados adequados, podem aprender habilidades sociais suficientes para se manter num emprego e articular seus processos mentais. E uma pequena parte deles tem notáveis capacidades de savants. Alguns especialistas acreditam que muitos grandes cientistas tinham síndrome de Asperger. Isso explicaria a natureza estranha, reclusa, de físicos como Isaac Newton e Paul Dirac (um dos fundadores da teoria quântica). Newton, em especial, era patologicamente incapaz de conversar sobre amenidades.

Tive o prazer de entrevistar um desses indivíduos, Daniel Tammet, autor do best-seller *Nascido em um dia azul*. Um dos únicos savants notáveis, ele é capaz de articular seus pensamentos em livros e entrevistas em rádio e televisão. Para alguém que tinha tanta dificuldade de relacionamento quando criança, ele tem hoje uma excelente capacidade de comunicação.

Daniel se distinguiu por conseguir um recorde mundial de memorização do Pi, um número fundamental na geometria. Ele foi capaz de memorizá-lo com 22.514 casas decimais. Perguntei como ele tinha se preparado para tal proeza hercúlea, e Daniel me respondeu que associa uma cor ou textura a cada número. Então, fiz a pergunta-chave: Se todo dígito tem uma cor ou textura, como você consegue recordar dezenas de milhares delas? Infelizmente, ele disse que não sabia. Apenas lhe ocorre. Os números têm sido sua vida desde criança, e simplesmente aparecem em sua mente, que é uma mistura constante de números e cores.

ASPERGER E O VALE DO SILÍCIO

Até aqui, esse debate pode parecer abstrato, sem nenhuma relação direta com nossa vida. Mas o impacto das pessoas com autismo leve e Asperger pode ser mais amplo do que se pensava, especialmente nos campos da alta tecnologia.

Na popular série de televisão *The Big Bang Theory*, vemos as palhaçadas de jovens cientistas, principalmente físicos nerds, numa procura desastrada por uma namorada. Em cada episódio, há um incidente engraçado que revela como eles são desinformados e patéticos nesse sentido.

Há uma suposição tácita nessa série de TV, de que a inteligência brilhante

deles é comparável apenas à sua esquisitice. E diz-se, meio de brincadeira, que entre os gurus da alta tecnologia no Vale do Silício há uma percentagem maior que a normal de pessoas que carecem de aptidão social. Entre as mulheres que frequentam universidades especializadas em engenharia, onde a proporção de homens é decididamente a favor delas, há um ditado: *"The odds are good, but the goods are odd."*^[3] Os cientistas decidiram investigar essa suspeita. A hipótese é de que pessoas com Asperger e outras formas amenas de autismo têm capacidades mentais ideais para certas áreas, como a indústria de tecnologia da informação. Cientistas da University College London examinaram 16 pessoas com diagnóstico de uma forma leve de autismo e as compararam com 16 indivíduos normais. Os dois grupos assistiram a slides contendo letras e números aleatórios, em padrões cada vez mais complexos.

Os resultados mostraram que os indivíduos com autismo tinham uma capacidade superior de se concentrar na tarefa. De fato, à medida que a tarefa ficava mais difícil, a diferença entre a capacidade intelectual dos grupos era cada vez maior, com o desempenho dos autistas significativamente melhor. Por outro lado, o teste mostrou que esses indivíduos se distraíam mais facilmente com ruídos externos e luzes piscantes do que o grupo controle.

A dra. Nilli Lavie diz: "Nosso estudo confirma a hipótese de que pessoas com autismo têm maior capacidade de percepção em comparação com a população normal. (...) Pessoas com autismo são capazes de perceber mais informações do que um adulto típico."

Certamente, isso não prova que todas as pessoas intelectualmente brilhantes têm alguma forma de Asperger. Mas indica, sim, que certas áreas que exigem maior capacidade de concentração intelectual podem ter uma proporção mais alta de indivíduos com Asperger.

VARREDURA CEREBRAL DE SAVANTS

O tema savantismo sempre foi rodeado de boatos e casos fantásticos. Mas, recentemente, essa área foi revirada de cabeça para baixo com o desenvolvimento da IRM e outras varreduras cerebrais.

O cérebro de Kim Peek, por exemplo, era incomum. A IRM mostrou que lhe faltava o corpo caloso ligando o cérebro esquerdo ao direito, o que provavelmente explica ele conseguir ler duas páginas ao mesmo tempo. Sua deficiência motora foi verificada numa deformação do cerebelo, a área que controla o equilíbrio. Infelizmente, a IRM não é capaz de revelar a origem exata de suas extraordinárias habilidades e memória fotográfica. Mas, de um modo geral, varreduras cerebrais mostraram que muitos com síndrome de savant

adquirida sofreram algum dano no lado esquerdo do cérebro.

Os estudos têm se concentrado particularmente nos córtices temporal anterior esquerdo e orbitofrontal. Alguns acreditam que talvez *todas* as habilidades dos savants (seja em autistas, adquiridos, ou aspergers) resultam de um dano nesse local muito específico do lobo temporal esquerdo. Essa área pode agir como um “censor” que elimina periodicamente lembranças irrelevantes. Mas quando ocorre um dano no hemisfério esquerdo, o hemisfério direito assume. O cérebro direito é muito mais preciso do que o esquerdo, que frequentemente distorce a realidade e cria fabulações. Na verdade, acredita-se que o cérebro direito passa a trabalhar muito mais quando o cérebro esquerdo é lesado, e o desenvolvimento das habilidades típicas dos savants é uma das consequências. Por exemplo: o cérebro direito é muito mais artístico do que o esquerdo. Normalmente, o lado esquerdo restringe esse talento, e o mantém sob controle. Mas se o cérebro esquerdo é lesado de determinada forma, pode liberar as habilidades artísticas latentes no cérebro direito, provocando uma explosão de talento artístico. Assim, a chave para liberar as habilidades savants poderia ser a atenuação do cérebro esquerdo para que não restrinjam mais os talentos naturais do cérebro direito. Isso costuma ser chamado de “lesão no esquerdo, compensação no direito”.

Em 1998, o dr. Bruce Miller, da Universidade da Califórnia em São Francisco, conduziu uma série de estudos, que parecem sustentar essa ideia. Ele e seus colaboradores estudaram cinco indivíduos normais que começaram a apresentar sinais de demência frontotemporal (DFT). À medida que a demência avançava, apareciam as capacidades savants. Quando a demência piorou, vários desses indivíduos começaram a apresentar talentos artísticos cada vez mais extraordinários, que nunca tinham se manifestado antes. Além disso, os talentos eram típicos do comportamento de savants. As habilidades eram visuais, não auditivas, e seus trabalhos artísticos, embora notáveis, não passavam de cópias, sem qualidades abstratas, simbólicas, nem originalidade. (Uma paciente até melhorou durante o estudo, mas em consequência disso seus talentos de savant diminuíram. Isso sugere uma relação próxima entre distúrbios do lobo temporal esquerdo e o surgimento de habilidades savants.)

A análise do dr. Miller aparentemente demonstrou que a degeneração dos córtices temporal anterior esquerdo e orbitofrontal provavelmente reduziram a inibição dos sistemas visuais no hemisfério direito, aumentando assim as habilidades artísticas. Mais uma vez, danos num local específico do hemisfério esquerdo forçaram o hemisfério direito a assumir o comando e se desenvolver.

Além dos savants, também foram realizadas varreduras de IRM em indivíduos com síndrome hipertímica, que também apresentam memória fotográfica. Não sofrem de autismo nem de distúrbios mentais, mas têm algumas características desses pacientes. Em todos os Estados Unidos, foram documentados apenas quatro casos de memória fotográfica desse tipo. Um deles

é o de Jill Price, administradora escolar em Los Angeles. Ela consegue lembrar exatamente o que estava fazendo em qualquer dia, inclusive décadas atrás. E se queixa de não conseguir apagar certos pensamentos. De fato, seu cérebro parece estar “travado no piloto automático”. Ela compara sua memória a contemplar o mundo numa tela dividida, onde o passado e o presente estão sempre disputando sua atenção.

Cientistas da Universidade da Califórnia em Irvine vêm fazendo varreduras no cérebro de Jill desde o ano 2000, e descobriram anomalias. Várias regiões são maiores que o normal, como o núcleo caudado (que trata da formação de hábitos) e o lobo temporal (que armazena fatos e números). Teorizou-se que essas duas áreas operam em conjunto para criar a memória fotográfica. Seu cérebro, portanto, é diferente do cérebro de savants que sofreram lesão ou algum dano no lobo temporal esquerdo. O motivo é desconhecido, mas indica outro caminho para se chegar a essas fantásticas capacidades mentais.

PODEMOS NOS TORNAR SAVANTS?

Tudo isso aponta para a instigante possibilidade de alguém ser capaz de desativar partes do cérebro esquerdo para aumentar a atividade do hemisfério direito, forçando-o a adquirir habilidades savants.

Lembramos que a estimulação magnética transcraniana (EMT) permite silenciar efetivamente partes do cérebro. Sendo assim, por que não podemos silenciar partes dos córtices temporal anterior esquerdo e orbitofrontal, usando a EMT, para nos tornarmos intencionalmente gênios savants?

Na verdade, essa ideia já foi testada. O dr. Allan Snyder, da Universidade de Sydney, na Austrália, ganhou as manchetes alguns anos atrás quando divulgou que, com aplicações de EMT em certa parte do cérebro esquerdo, os indivíduos de seu experimento conseguiram subitamente realizar façanhas de savants. Ao direcionar ondas magnéticas de baixa frequência para o hemisfério esquerdo, pode-se, em princípio, desligar essa região dominante, e o hemisfério direito passa a dominar. O dr. Snyder e colegas fizeram um experimento com 11 voluntários homens. Aplicaram a EMT na região frontotemporal esquerda desses homens durante testes com leitura e desenho. Os voluntários não desenvolveram habilidades savants, mas dois deles apresentaram melhora significativa na capacidade de revisar textos e detectar palavras duplicadas. Em outro experimento, o dr. R. L. Young e colegas aplicaram uma bateria de testes psicológicos em 17 indivíduos. Os testes eram específicos para avaliar habilidades savants. (Esse tipo de teste analisa a capacidade de memorização, de manipulação de números e datas, além de criação artística e desempenho

musical.) Cinco indivíduos apresentaram progressos nas habilidades típicas de savants depois do tratamento com EMT.

O dr. Michael Sweeney observou: “Quando aplicada aos lobos pré-frontais, a EMT mostrou que aumentam a rapidez e a agilidade dos processos cognitivos. Os disparos da EMT agem como aplicações localizadas de cafeína, mas ninguém sabe ao certo como os magnetos realmente operam.” Esses experimentos indicam, mas não provam, que silenciar uma parte da região frontotemporal esquerda pode desenvolver algumas habilidades. Mas estão longe de ser as dos savants, e devemos também salientar que outros grupos revisaram esses experimentos e os resultados não foram conclusivos. É preciso realizar novas experimentações; ainda é muito cedo para dar uma palavra final concordando ou discordando.

As sondas de EMT são os instrumentos mais fáceis e convenientes para essa finalidade, pois podem silenciar intencionalmente várias partes selecionadas do cérebro, sem acidentes traumáticos nem danos cerebrais. Mas é preciso ressaltar que a sonda de EMT ainda é muito primitiva, silenciando milhões de neurônios de uma vez. Os campos magnéticos, diferentemente das sondas elétricas, não têm muita precisão, mas se estendem por vários centímetros. Sabemos que os córtices temporal anterior esquerdo e orbitofrontal de savants apresentam lesões que provavelmente são responsáveis, pelo menos em parte, por suas capacidades singulares, mas talvez a área específica a ser atenuada seja uma sub-região ainda menor. Assim, uma aplicação de EMT pode desativar algumas áreas que precisam ficar intactas para produzir habilidades de savants.

No futuro, com as sondas de EMT poderemos delimitar ainda mais a região do cérebro envolvida no desenvolvimento das habilidades dos savants. Quando essa região for identificada, o passo seguinte será usar sondas elétricas de alta precisão, como as usadas na estimulação cerebral profunda, para atenuar tais áreas com maior exatidão. Então bastará apertar um botão para que essas sondas silenciem uma minúscula porção do cérebro, a fim de despertar as habilidades savants.

ESQUECENDO DE ESQUECER E MEMÓRIA FOTOGRÁFICA

Embora as habilidades savants possam ser ativadas por uma lesão no cérebro esquerdo (exigindo uma compensação pelo cérebro direito), isso ainda não explica exatamente como o cérebro direito consegue um desempenho dessa magnitude em termos de memorização. Qual é o mecanismo neural que faz surgir a memória fotográfica? A resposta a essa pergunta talvez determine se poderemos nos tornar savants.

Até recentemente, pensava-se que a memória fotográfica era fruto de uma capacidade especial de certos cérebros. Se assim fosse, dificilmente uma pessoa comum poderia adquirir essa habilidade de memorização, que seria privilégio de cérebros excepcionais. Mas, em 2012, um novo estudo mostrou que a verdade pode estar no extremo oposto.

A explicação para a memória fotográfica não deve ser uma capacidade de cérebros privilegiados; pelo contrário, deve ser sua incapacidade de esquecer. Se isso for verdade, talvez a memória fotográfica não seja algo tão misterioso assim.

O novo estudo foi realizado por cientistas do Scripps Research Institute na Flórida, que trabalhavam com moscas-das-frutas. Eles descobriram um modo interessante de aprendizado das moscas, que pode derrubar a ideia estabelecida sobre formação de lembranças e esquecimento. As moscas foram expostas a diferentes odores e receberam reforço positivo (com comida) ou reforço negativo (com choques elétricos).

Os cientistas sabem que o neurotransmissor dopamina é importante para a formação de memória. Para surpresa de todos, descobriram que a dopamina regula ativamente *tanto* a formação *como* o esquecimento de novas lembranças. No processo de criação de memória, o receptor dCA1 foi ativado. Em contraste, o esquecimento foi iniciado pela ativação do receptor DAMB.

Anteriormente pensava-se que o esquecimento era simplesmente a degradação das lembranças com o passar do tempo, um processo passivo que ocorre por si só. Esse novo estudo mostra que esquecer é um processo ativo, que exige a intervenção da dopamina.

Para provar a descoberta, eles mostraram que, interferindo na ação dos receptores dCA1 e DAMB, podiam aumentar ou diminuir propositalmente a capacidade de lembrança ou esquecimento das moscas-das-frutas. Observaram que uma mutação no receptor dCA1, por exemplo, prejudica a capacidade de lembrar, e uma mutação no receptor DAMB diminui a capacidade de esquecer.

Os pesquisadores especulam se esse efeito, por sua vez, pode ser parcialmente responsável pelas habilidades dos savants. Talvez eles tenham uma deficiência na capacidade de esquecer. Um dos estudantes que participou da pesquisa, Jacob Berry, diz: “Os savants têm uma alta capacidade de memória. Mas talvez não seja a memória que lhes dá essa capacidade; talvez seja um mau mecanismo de esquecimento. Isso pode vir a ser uma estratégia para o desenvolvimento de medicamentos que melhorem a memória e a cognição – que tal, medicamentos inibidores do esquecimento para melhorar a cognição?”

Supondo-se que esse resultado se aplique da mesma forma a humanos, pode estimular os cientistas a desenvolver novos medicamentos e neurotransmissores capazes de enfraquecer o processo de esquecimento. Neutralizando tal processo, deve ser possível ativar seletivamente a memória fotográfica quando for preciso. Dessa maneira talvez se consiga eliminar o excesso de informações irrelevantes

e inúteis, que atrapalham o pensamento de pessoas com síndrome de savant.

Também é animadora a possibilidade de que o projeto Brain, apoiado pelo governo de Obama, possa identificar os caminhos específicos envolvidos na síndrome de savant adquirida. O desenvolvimento dos campos magnéticos transcranianos ainda é muito incipiente para isolar os poucos neurônios que possam estar envolvidos. Mas, com o uso de nanossondas e das últimas tecnologias de varreduras, o projeto Brain poderá isolar com exatidão os percursos neurais que possibilitam a memória fotográfica e as habilidades computacionais, artísticas e musicais fora do comum. Bilhões de dólares destinados a essa pesquisa serão direcionados para a identificação dos caminhos neurais específicos envolvidos em doenças mentais e outros problemas cerebrais, e o segredo das habilidades dos savants poderá ser revelado. Então, talvez seja possível que indivíduos normais se tornem savants também. Isso já aconteceu muitas vezes no passado em consequência de acidentes variados. No futuro, pode vir a ser um procedimento médico de alta precisão. O tempo dirá.

Até o momento, os métodos analisados aqui não alteram a natureza do cérebro, nem do resto do corpo. Há esperança de que, com o uso de campos magnéticos, sejamos capazes de liberar o potencial existente no cérebro em estado de latência. A filosofia por trás dessa ideia é que somos todos savants à espera de uma manifestação, e bastaria uma pequena alteração em nossos circuitos neurais para liberar talentos ocultos.

Uma outra técnica seria alterar diretamente o cérebro e os genes, utilizando o que há de mais moderno na ciência do cérebro e também na genética. Um método promissor é o uso de células-tronco.

CÉLULAS-TRONCO PARA O CÉREBRO

Durante muitas décadas, imperou o dogma de que as células cerebrais não se regeneram. Parecia impossível renovar células velhas, quase mortas, ou criar outras células para potencializar nossas capacidades. Mas tudo mudou em 1998. Naquele ano, descobriu-se que era possível encontrar células-tronco no hipocampo, no bulbo olfativo, e no núcleo caudado de adultos. Resumindo, as células-tronco são a “mãe de todas as células”. As células-tronco de embriões, por exemplo, podem se desenvolver rapidamente, transformando-se em qualquer outra célula.

Embora cada uma de nossas células contenha todo o material genético necessário para gerar um ser humano, somente as células-tronco embrionárias podem se diferenciar, transformando-se em qualquer outro tipo de célula.

Células-tronco de adultos perderam essa capacidade camaleônica, mas ainda

podem se reproduzir e substituir células velhas e quase mortas. Em se tratando de melhorar a memória, a atenção se volta para as células-tronco no hipocampo de adultos. Percebeu-se que milhares de células novas nascem naturalmente todos os dias no hipocampo, mas muitas morrem logo depois. Entretanto, viu-se que ratos que aprenderam novas habilidades retiveram uma maior quantidade de células novas. Uma combinação de exercícios e substâncias químicas estimulantes do humor também pode aumentar muito a taxa de sobrevivência das células novas do hipocampo. E já se constatou que o estresse, pelo contrário, acelera a morte de neurônios novos.

Em 2007, houve um grande avanço quando cientistas em Wisconsin e no Japão conseguiram colher células normais da pele humana, reprogramar seus genes e torná-las células-tronco. Há esperança de que as células-tronco, encontradas naturalmente ou convertidas pela engenharia genética, possam um dia ser injetadas no cérebro de pacientes com Alzheimer para substituir as células quase mortas. Essas novas células cerebrais, como ainda não têm as conexões apropriadas, não se integram à arquitetura neural. Isso significa que seria necessário reaprender certas habilidades para incorporar os novos neurônios.

Naturalmente, o estudo de células-tronco é uma das áreas mais produtivas das pesquisas sobre o cérebro. “As pesquisas de células-tronco e a medicina regenerativa estão numa fase extremamente animadora. Estamos adquirindo conhecimento com muita rapidez, e estão surgindo diversas empresas que fazem experimentos clínicos em várias áreas”, diz o sueco Jonas Frisén, do Karolinska Institute.

A GENÉTICA DA INTELIGÊNCIA

Além das células-tronco, outra área de investigação se abre com o isolamento dos genes responsáveis pela inteligência humana. Os biólogos afirmam que somos 98,5% geneticamente idênticos aos chimpanzés. No entanto, vivemos duas vezes mais e nos últimos seis milhões de anos tivemos um progresso excepcional em termos de capacidade intelectual. Portanto, os genes responsáveis pelo desenvolvimento do cérebro humano devem estar em um grupo reduzido de genes. Dentro de poucos anos os cientistas terão um mapa completo de todas as diferenças genéticas, e o segredo da longevidade e da inteligência superior humana poderá ser encontrado. Os cientistas se concentraram em alguns genes que, possivelmente, impulsionaram a evolução do cérebro humano.

A pista para revelar o segredo da inteligência talvez esteja no conhecimento sobre nossos ancestrais símiescos. Isso levanta outra questão: Essa pesquisa tornará possível um *Planeta dos macacos*?

Nessa longa série de filmes, uma guerra nuclear destrói a civilização moderna. A humanidade é reduzida à barbárie, mas a radiação de algum modo acelera a evolução dos primatas, que se tornam a espécie dominante no planeta. Eles criam uma civilização avançada, enquanto os humanos são reduzidos a selvagens imundos e malcheirosos perambulando seminus pelas florestas. Os humanos são, no máximo, animais de zoológico. O jogo virou contra os humanos, e os macacos é que ficam nos encarando do lado de fora das jaulas.

Em um dos últimos filmes da série, *Planeta dos macacos – A origem*, os cientistas estão procurando a cura do mal de Alzheimer, e esbarram num vírus que por acaso provoca o aumento da inteligência dos chimpanzés. Infelizmente, um dos macacos que ficou mais inteligente é tratado com crueldade ao ser levado para um abrigo de primatas. Usando a nova inteligência, o macaco se liberta, contagia todos os animais do laboratório com o vírus que expande a inteligência, e abre as jaulas de todos eles. Pouco depois, uma multidão de macacos inteligentes aparece gritando, toma o controle da ponte Golden Gate, domina todo o local e intimida a polícia. Após um confronto espetacular e angustiante com as autoridades, o filme termina com os macacos se refugiando pacificamente numa floresta de sequoias ao norte da ponte.

Essa história é realista? Em curto prazo, não, mas não pode ser descartada, pois nos próximos anos os cientistas devem catalogar todas as mudanças genéticas que criaram o *Homo sapiens*. Mas temos que solucionar muitos outros mistérios antes de chegarmos aos macacos inteligentes.

Uma cientista fascinada, não pela ficção científica, mas pela genética do que nos faz “humanos”, é a dra. Katherine Pollard, especialista numa área chamada “bioinformática”, que mal existia uma década atrás. Nesse campo da biologia, em vez de abrir animais para entender como são compostos, os pesquisadores usam o enorme poder dos computadores para analisar matematicamente os genes do corpo deles. Ela está à frente das pesquisas sobre os genes que definem a essência do que nos separa dos macacos. Em 2003, recém-graduada doutora pela Universidade da Califórnia em Berkeley, ela teve chance de levar adiante sua pesquisa.

“Agarrei a oportunidade de participar da equipe internacional que estava identificando a sequência de bases de DNA, ou ‘letras’, no genoma do chimpanzé comum”, ela recorda. Sua meta era clara. Ela sabia que somente 15 milhões de pares de bases, ou “letras” que compõem nosso genoma (dentre três bilhões de pares de base) nos separam dos chimpanzés, nossos vizinhos genéticos mais próximos. (Cada “letra” de nosso código genético se refere a um dos quatro ácidos nucleicos, simbolizados por A, T, C e G. Assim, nosso genoma consiste em três bilhões de letras, em sequências do tipo ATTCCAGGG...)

“Eu estava determinada a descobri-las”, disse a doutora.

Isolar esses genes pode ter enormes implicações no futuro. Quando

conhecermos os genes que deram origem ao *Homo sapiens*, será possível determinar como os humanos evoluíram. O segredo da inteligência pode estar nesses genes. Pode até ser possível acelerar o processo da evolução e desenvolver nossa inteligência. Mas, ainda assim, 15 milhões de pares é um número gigantesco para ser analisado. Como achar um punhado de agulhas genéticas nesse palheiro genético?

A dra. Pollard sabia que a maior parte do nosso genoma é composto por “DNA lixo”, que não contém genes e praticamente não se alterou com a evolução. Esse DNA lixo sofre mutações lentamente (aproximadamente 1% muda a cada 4 milhões de anos). Como nossa diferença para os chimpanzés é 1,5% do DNA, isso significa que provavelmente nos separamos dos chimpanzés há cerca de 6 milhões de anos. Portanto, existe um “relógio molecular” em cada uma de nossas células. E, como a evolução acelera a taxa de mutação, uma análise de onde ocorreu essa aceleração permite dizer quais genes estão conduzindo a evolução.

A dra. Pollard ponderou que, se pudesse criar um programa de computador para descobrir onde se localiza, em nosso genoma, a maior parte dessas mudanças, ela poderia isolar com exatidão os genes que deram origem ao *Homo sapiens*. Após anos de muito trabalho e aperfeiçoamento, ela finalmente instalou seu programa nos gigantescos computadores da Universidade da Califórnia em Santa Cruz. E ficou aguardando ansiosamente os resultados.

Quando finalmente conseguiu imprimir o resultado, ela encontrou o que procurava: 201 regiões do nosso genoma mostravam uma mudança acelerada. Mas foi a primeira da lista que chamou sua atenção.

“Com meu mentor David Haussler olhando por cima do meu ombro, vi no topo da lista uma sequência de 118 bases que, juntas, ficaram conhecidas como região acelerada humana 1 (HAR1)”, disse. Ela ficou extasiada. Bingo! “Acertamos na loteria”, escreveu. Era a realização de um sonho.

Diante dela estava uma área do nosso genoma contendo apenas 118 pares de bases, com a maior divergência de mutações nos separando dos macacos. Nesses pares de bases, apenas 18 mutações foram alteradas desde que nos tornamos humanos. Sua notável descoberta mostrou que umas poucas mutações terão sido responsáveis por nos desatolar dos pântanos do passado genético.

Em seguida, ela e seus colegas tentaram decifrar a natureza exata dessa misteriosa região chamada HAR1, e viram que permaneceu espantosamente estável por milhões de anos de evolução. Os primatas se separaram das galinhas há cerca de 300 milhões de anos e apenas dois pares de bases diferenciam chimpanzés de galinhas. Portanto, HAR1 tinha permanecido praticamente imutável durante centenas de milhões de anos, com apenas duas mudanças, nas letras G e C. No entanto, em apenas seis milhões de anos, HAR1 teve 18 mutações, o que representa uma enorme aceleração em nossa evolução.

O mais enigmático, porém, era o papel da HAR1 no controle da configuração geral do córtex cerebral, famoso pela aparência enrugada. Um defeito na região HAR1 provoca um distúrbio chamado “liscencefalia”, ou “cérebro liso”, que causa dobras incorretas no cérebro. (Defeitos nessa região são também ligados à esquizofrenia.) Além do tamanho grande do nosso córtex cerebral, outra de suas características é ser muito enrugado e convoluto, o que aumenta muito a área de superfície e, portanto, a capacidade computacional. O trabalho da dra. Pollard mostrou que a variação de apenas 18 letras de nosso genoma era parcialmente responsável por uma das mais importantes e determinantes mudanças genéticas na história humana, amplamente expandindo nossa inteligência. (Lembrem que o cérebro de um dos maiores matemáticos da história, Carl Friedrich Gauss, preservado após sua morte, exibiu um enrugamento incomum.)

A lista da dra. Pollard foi ainda mais além, identificando outras centenas de áreas que também acusavam mudança acelerada, algumas das quais já eram conhecidas. A FOX2, por exemplo, é crucial para o desenvolvimento da fala, outra característica fundamental dos humanos. (Indivíduos com o gene FOX2 defeituoso têm dificuldade em fazer os movimentos faciais necessários à fala.) Outra região, chamada HAR2, dá aos nossos dedos a destreza para manusear instrumentos delicados.

Além disso, desde que o genoma do homem de Neandertal foi sequenciado, tornou-se possível comparar nossa constituição genética com a de uma espécie ainda mais próxima que os chimpanzés. Ao analisar o gene FOX2 em homens de Neandertal, os cientistas viram que temos o mesmo gene que eles. Isso significa a possibilidade de que o homem de Neandertal podia vocalizar e criar falas como nós.

Outro gene crucial é chamado ASPM, supostamente responsável pelo crescimento explosivo de nossa capacidade cerebral. Alguns cientistas acreditam que esse e outros genes podem revelar por que os humanos se tornaram inteligentes e os macacos não. (Pessoas com uma versão defeituosa do gene ASPM geralmente sofrem de microcefalia, uma forma grave de retardo mental, porque têm o cérebro muito pequeno, do tamanho do cérebro de nossos ancestrais australopitecos.)

Os cientistas rastream o gene ASPM e descobriram que sofreu cerca de 15 mutações nos últimos cinco ou seis milhões de anos, desde que nos separamos dos chimpanzés. Mutações mais recentes nesses genes parecem estar correlacionadas a marcos importantes em nossa evolução. Por exemplo: ocorreu uma mutação 100 mil anos atrás, quando os humanos modernos surgiram na África, com aparência indistinguível da nossa. E a última mutação ocorreu há 5.800 anos, o que coincide com a introdução da linguagem escrita e da agricultura.

Como essas mutações coincidem com períodos de rápido crescimento do

intelecto, é instigante especular que o ASPM esteja entre os poucos genes responsáveis pela nossa maior inteligência. Se isso for verdade, talvez possamos determinar se esses genes ainda estão ativos hoje, e se continuarão a moldar a evolução humana no futuro.

Toda essa pesquisa levanta uma pergunta: A manipulação de alguns genes pode expandir nossa inteligência?

Muito possivelmente.

Os cientistas estão avançando rapidamente na direção do mecanismo exato pelo qual esses genes deram origem à inteligência. Particularmente, regiões genéticas e genes como o HAR1 e o ASPM podem ajudar a solucionar um mistério do cérebro: se existem aproximadamente 23 mil genes em nosso genoma, como é possível controlarem as conexões entre 100 bilhões de neurônios, totalizando um quatrilhão (um com quinze zeros) de conexões? Parece matematicamente impossível. O genoma humano é cerca de um trilhão de vezes menor do que precisaria ser para codificar todas as nossas conexões neurais. Assim, nossa própria existência parece ser uma impossibilidade matemática.

A resposta pode estar no fato de que a natureza pega inúmeros atalhos para criar o cérebro. Primeiro, muitos neurônios têm conexões aleatórias, portanto não há necessidade de um esquema detalhado. Isso significa que as regiões conectadas aleatoriamente se organizam depois que o bebê nasce e começa a interagir com o ambiente.

Segundo, a natureza usa módulos que se repetem diversas vezes. Quando a natureza descobre algo útil, quase sempre o repete. Isso pode explicar como apenas algumas mudanças genéticas são responsáveis por quase todo o desenvolvimento da nossa inteligência nos últimos seis milhões de anos.

Nesse caso, o tamanho é importante, sim. Se alterarmos o ASPM e alguns outros genes, o cérebro pode ficar maior e mais complexo, tornando possível expandir nossa inteligência. (Aumentar o tamanho do cérebro não basta, dado que a maneira como o cérebro é organizado também importa. Mas aumentar a massa cinzenta é uma pré-condição necessária para o aumento da inteligência.)

MACACOS, GENES E GÊNIOS

A pesquisa da dra. Pollard se concentrou em áreas do genoma que compartilhamos com os chimpanzés e que sofreram mutação. É possível também que existam áreas do genoma encontradas apenas em humanos, independentes do genoma dos macacos. Um gene desses foi descoberto recentemente, em novembro de 2012. Cientistas liderados por uma equipe da Universidade de Edimburgo isolaram o gene RIM-941, o único descoberto até

agora que é exclusivo do *Homo sapiens*; não é encontrado em primatas. Geneticistas também mostraram que esse gene surgiu entre seis e um milhões de anos atrás (ou seja, apenas depois que os humanos e os chimpanzés se separaram).

Infelizmente, essa descoberta também ecoou em publicações e blogs científicos, com manchetes equivocadas borbulhando na internet. Artigos precipitados anunciavam a descoberta de um único gene que, em princípio, podia tornar os chimpanzés inteligentes. A essência de nossa “humanidade” tinha finalmente sido isolada no nível genético, diziam as manchetes.

Cientistas conceituados logo se manifestaram, tentando acalmar os ânimos. É provável que uma série de genes, em ação conjunta e muito complexa, seja responsável pela inteligência humana. Nenhum gene isolado pode fazer um macaco ter a inteligência de um humano de um dia para o outro, disseram.

Apesar do grande exagero das manchetes, isso de fato levantou uma questão muito séria: O que há de realidade no *Planeta dos macacos*?

Existe aí uma série de complicações. Se os genes HAR1 e ASPM forem aprimorados a ponto de expandir repentinamente a estrutura do cérebro dos chimpanzés, uma série de outros genes teria que ser modificada também. Primeiro, seria preciso fortalecer os músculos do pescoço e aumentar o tamanho do corpo dos chimpanzés para que suportem uma cabeça maior. Mas um cérebro grande seria inútil se não conseguisse controlar os dedos e manusear instrumentos. Assim, o gene HAR2 também teria que ser alterado para aumentar a destreza. E como os chimpanzés geralmente andam com o apoio das mãos, outros genes teriam que ser alterados para que a coluna vertebral sustentasse uma postura ereta, liberando as mãos. A inteligência seria inútil se os chimpanzés não pudessem se comunicar pela fala com outros indivíduos da espécie. Portanto, o gene FOX2 também teria que sofrer mutação para tornar possível uma fala semelhante à humana. Por fim, para criar uma espécie de macacos inteligentes, seria preciso alterar o canal vaginal, que não tem largura suficiente para permitir a passagem de um crânio maior. Seria necessário fazer cesariana para tirar o filhote, ou alterar geneticamente o canal vaginal das fêmeas.

Depois de todas essas adaptações genéticas, teremos uma criatura muito parecida conosco. Em outras palavras, pode ser anatomicamente impossível criar macacos inteligentes como vemos no cinema sem criar mutações que os deixem muito parecidos com seres humanos.

Portanto, é claro que criar macacos inteligentes não é tão fácil. Os macacos inteligentes dos filmes de Hollywood são humanos fantasiados ou imagens geradas por computador. Assim, todas essas questões são convenientemente varridas para baixo do tapete. Mas se os cientistas pudessem realmente usar a terapia de genes para criar macacos inteligentes, estes seriam muito parecidos conosco, com mãos adequadas para usar instrumentos, cordas vocais para a fala,

coluna vertebral para sustentar a postura ereta e músculos fortes no pescoço para apoiar uma cabeça grande, tal como nós.

Tudo isso levanta questões éticas. Embora a sociedade permita estudos genéticos com macacos, pode não tolerar a manipulação de seres inteligentes que sintam dor e angústia. Afinal, essas criaturas seriam inteligentes e articuladas o suficiente para se queixar de sua situação e seu destino, e a opinião delas seria ouvida pela sociedade.

Como se sabe, essa área da bioética é tão nova que é totalmente inexplorada. Tal tecnologia ainda não está pronta, mas nas próximas décadas, quando tivermos identificado todos os genes e suas funções que nos separam dos macacos, o tratamento desses animais modificados pode se tornar um tema crítico.

Vemos, portanto, que é apenas uma questão de tempo até que todas as pequenas diferenças genéticas entre nós e os chimpanzés sejam minuciosamente sequenciadas, analisadas e interpretadas. Mas resta uma questão ainda mais profunda: Quais foram as forças evolutivas que nos deram essa herança genética depois que nos separamos dos macacos? Por que genes como o ASPM, HAR1 e FOX2 vieram a se desenvolver? Em outras palavras, a genética nos dá a capacidade de entender como nos tornamos inteligentes, mas não explica por que isso aconteceu.

Se pudermos entender esse problema, talvez consigamos pistas para saber como poderemos evoluir no futuro. Isso nos leva ao cerne deste debate: Qual é a origem da inteligência?

A ORIGEM DA INTELIGÊNCIA

Muitas teorias, remontando a Charles Darwin, tentaram explicar por que os humanos desenvolveram uma inteligência maior.

Segundo uma das teorias, a evolução do cérebro humano ocorreu em estágios, tendo a primeira fase se iniciado na mudança climática da África. À medida que o clima esfriava, as florestas diminuía, expulsando nossos ancestrais para planícies abertas, savanas, onde ficavam mais expostos aos predadores e à natureza. Para sobreviver nesse ambiente mais hostil, foram obrigados a caçar e a andar eretos, o que liberou suas mãos e gerou a oposição dos polegares para manusear ferramentas. Isso, por sua vez, privilegiou os cérebros maiores, que coordenavam a fabricação dessas ferramentas. Segundo tal teoria, o homem antigo não fazia ferramentas – “as ferramentas fizeram o homem”.

Nossos ancestrais não descobriram as ferramentas de repente e ficaram inteligentes. Aconteceu o oposto. Os humanos que descobriram como usar instrumentos sobreviveram nas planícies, e os que não usavam foram se

extinguindo. Os humanos que sobreviveram e se desenvolveram nas planícies foram aqueles que, através de mutações, se adaptaram cada vez mais ao uso e manufatura de ferramentas, o que exigia um cérebro cada vez maior.

Outra teoria destaca nossa natureza social, coletiva. Os humanos conseguem coordenar facilmente o comportamento de mais de 100 indivíduos envolvidos numa situação de caça, agricultura, guerra e construção, grupos muito maiores do que os encontrados entre os outros primatas, o que deu aos humanos uma vantagem sobre os demais animais. Segundo essa teoria, é preciso um cérebro maior para ser capaz de distribuir tarefas e controlar tantos indivíduos. O outro lado da teoria indica que é necessário um cérebro maior para conspirar, tramar, enganar e manipular os outros seres inteligentes da tribo. Os indivíduos capazes de entender os motivos dos outros e explorá-los têm uma vantagem sobre os desprovidos dessa capacidade. É a teoria maquiavélica da inteligência.

Uma terceira teoria sustenta que o desenvolvimento da linguagem, que veio depois, acelerou o surgimento da inteligência. A linguagem trouxe o pensamento abstrato e a capacidade de planejar, organizar a sociedade, mapear etc. Os humanos têm um vocabulário mais extenso que qualquer outro animal, com dezenas de milhares de palavras conhecidas por uma pessoa comum. Com a linguagem, os humanos podiam coordenar e monitorar as atividades de muitos indivíduos, além de lidar com ideias e conceitos abstratos. A linguagem significava poder conduzir grupos de pessoas numa caçada, o que é uma grande vantagem no abate de um mamute grandalhão. Podiam dizer aos outros onde havia caça com fartura, e onde o perigo estava à espreita.

Uma outra teoria, da seleção sexual, defende que as fêmeas preferem se acasalar com machos inteligentes. No reino animal, como numa alcateia, por exemplo, o macho alfa mantém o grupo unido pela força bruta. Qualquer lobo que o desafiar tem que ser escorraçado a dentadas e patadas. Mas, milhões de anos atrás, à medida que os humanos ficavam cada vez mais inteligentes, a força bruta já não era suficiente para manter a união da tribo. Os mais astutos e inteligentes sabiam armar emboscadas, mentir, trapacear ou formar facções dentro da tribo para derrubar o macho alfa. Assim, o macho alfa das gerações seguintes já não era obrigatoriamente o mais forte. Com o tempo, o mais esperto e inteligente passou a ser o líder. Provavelmente, essa é a razão pela qual as fêmeas escolhem machos espertos (não necessariamente espertos do tipo nerd, mas que “sabem tudo”). A seleção sexual, por sua vez, acelerou nossa evolução na direção da inteligência. Nesse caso, o motor da expansão do nosso cérebro foram mulheres que escolheram homens capazes de formular estratégias, de assumir a liderança da tribo e superar outros homens em inteligência, o que requer um cérebro grande.

São poucas as teorias sobre a origem da inteligência, e todas têm prós e contras. O tema em comum parece ser a capacidade de simular o futuro. Por

exemplo: o objetivo do líder é escolher o melhor caminho para a tribo no futuro. Isso significa que o líder precisa entender as intenções dos outros a fim de planejar estratégias. Portanto, simular o futuro foi talvez uma das forças propulsoras da evolução para nosso cérebro atingir o tamanho e a inteligência atuais. E quem pode simular melhor o futuro é quem sabe tramar, conspirar, entender a mente dos outros membros da tribo, e estar sempre preparado para se defender de ataques.

Da mesma forma, a linguagem permite simular o futuro. Os animais possuem uma linguagem rudimentar, basicamente no tempo presente. Essa linguagem pode alertar sobre um perigo imediato, como um predador escondido na moita, mas, ao que parece, não possui tempo passado nem futuro. Os animais não conjugam verbos. Assim, talvez a capacidade de expressar os tempos passado e futuro tenha sido uma conquista determinante para o desenvolvimento da inteligência.

O dr. Daniel Gilbert, psicólogo de Harvard, escreve: “Nas primeiras centenas de milhões de anos, desde seu aparecimento no planeta, nosso cérebro ficou preso num eterno presente, e muitos ainda estão. Mas não o seu, nem o meu, porque dois ou três milhões de anos atrás nossos ancestrais iniciaram a grande escapada do aqui e agora.”

O FUTURO DA EVOLUÇÃO

Até aqui vimos resultados curiosos, indicando que é possível desenvolver a memória e a inteligência, mas isso se deve, em grande parte, à maior eficiência e maximização da capacidade natural do cérebro. Vários métodos estão sendo estudados, como medicamentos, genes e aparelhos (o de EMT, por exemplo) que podem aumentar as capacidades dos neurônios.

A ideia de alterar o tamanho do cérebro e a capacidade dos macacos é uma possibilidade, sim, embora seja difícil. A terapia de genes, nessa escala, ainda está a muitas décadas de distância. Mas isso levanta outras questões complicadas: Até onde isso pode chegar? Podemos estender indefinidamente a inteligência de um organismo? Ou as leis da física impõem um limite à modificação do cérebro?

A resposta para a última pergunta é sim. As leis da física impõem um limite ao que pode ser feito em termos de modificação genética. Há restrições sim. Para conhecer esse limite, em primeiro lugar, é necessário examinar se a evolução ainda está desenvolvendo a inteligência humana, e depois analisar o que pode ser feito para acelerar esse processo natural.

Na cultura popular há a ideia de que no futuro a evolução nos dará um cérebro enorme e um corpo pequeno e sem pelos. Coincidentemente, os alienígenas que

vêm do espaço, aos quais se atribui uma inteligência superior, são geralmente retratados assim. Qualquer loja de brinquedos hoje em dia expõe a mesma figura extraterrestre com imensos olhos de inseto, cabeça enorme e pele verde.

Na verdade, há indícios de que o grosso da evolução humana (isto é, a forma corporal básica e a inteligência) deu uma parada. Vários fatores sustentam tal suposição. Primeiro, como somos mamíferos bípedes que andam eretos, há limitações do tamanho máximo do crânio de um bebê para passar pelo canal vaginal. Segundo, o avanço da tecnologia moderna removeu muitas das pressões evolutivas severas enfrentadas por nossos ancestrais.

Entretanto, a evolução em termos genéticos e moleculares continua sempre. Embora seja difícil ver a olho nu, há evidências de que a bioquímica humana mudou para se adaptar aos percalços do ambiente, como, por exemplo, resistir à malária nos climas tropicais. Há pouco tempo, quando aprenderam a domesticar o gado e a beber o leite da vaca, os humanos desenvolveram enzimas que digerem a lactose. Ocorreram mutações à medida que os humanos se adaptavam à dieta criada pela revolução agrícola. Até hoje as pessoas ainda escolhem os mais adaptados e saudáveis como parceiros, e assim a evolução continua a eliminar genes inadequados nesse aspecto. Nenhuma dessas mutações, porém, mudou a constituição básica do nosso corpo, nem aumentou o tamanho do nosso cérebro.

Até certo ponto, a tecnologia moderna também influencia na evolução. Por exemplo, não existe mais uma seleção eliminando os míopes, pois todo mundo hoje pode usar óculos ou lentes de contato.

A FÍSICA DO CÉREBRO

Do ponto de vista evolutivo e biológico, a evolução já não seleciona os mais inteligentes, pelo menos não tão rapidamente como ocorria milhares de anos atrás.

Há também indicações, pelas leis da física, de que atingimos o limite natural máximo da inteligência, de modo que qualquer aumento deverá vir de meios externos. Os físicos que estudaram a neurologia do cérebro concluem que há desvantagens que nos impedem de ter mais inteligência. Cada vez que almejamos um cérebro maior, mais complexo ou mais denso, somos confrontados por tais questões.

O primeiro princípio da física aplicável ao cérebro é a conservação de matéria e energia, ou seja, a lei que estabelece que a quantidade total de matéria e energia num sistema permanece constante. A fim de realizar sua incrível ginástica mental, o cérebro precisa conservar energia e, para isso, toma vários

atalhos. Como vimos no capítulo 1, o que enxergamos com os olhos é na verdade recomposto com truques para poupar energia. Como fazer uma análise aprofundada de cada crise gastaria muito tempo e energia, o cérebro a economiza com julgamentos rápidos na forma de emoções. O esquecimento é uma outra maneira de economizar energia. O cérebro consciente só tem acesso a uma pequena porção das lembranças que têm impacto sobre ele.

Portanto, a pergunta é: O aumento do tamanho do cérebro ou da densidade dos neurônios nos daria mais inteligência?

Provavelmente, não. “Os neurônios da massa cinzenta cortical funcionam com axônios muito próximos do limite físico”, diz o dr. Simon Laughlin, da Universidade de Cambridge. Há várias maneiras de expandir a inteligência usando as leis da física, mas todas elas apresentam problemas:

- Pode-se aumentar o tamanho do cérebro e o comprimento dos neurônios. O problema aqui é que o cérebro consome mais energia. Esse processo gera mais calor, o que é prejudicial à nossa sobrevivência. Se o cérebro usa mais energia, fica mais quente, e se a temperatura do corpo aumenta demais, provoca danos nos tecidos. (As reações químicas do corpo humano e nosso metabolismo exigem que a temperatura se mantenha dentro de uma determinada faixa.) Além disso, neurônios mais longos implicam mais tempo para os sinais atravessarem o cérebro, o que retarda o processo de pensamento.
- Pode-se armazenar mais neurônios num mesmo espaço, tornando-os mais finos. Mas se os neurônios ficam cada vez mais finos, as complexas reações químicas e elétricas que ocorrem dentro dos axônios começam a falhar, e cada vez mais eles não disparam como deviam. Douglas Fox, em um artigo da *Scientific American*, diz que: “Essa pode ser considerada a mãe de todas as limitações: as proteínas que os neurônios usam para gerar pulsos elétricos, chamadas canais iônicos, são inerentemente instáveis.”
- Pode-se aumentar a velocidade do sinal, tornando os neurônios mais espessos. Mas isso também aumenta o consumo de energia e gera mais calor. E também aumenta o tamanho do cérebro, fazendo com que os sinais levem mais tempo para atingirem seu destino.
- Podem-se acrescentar mais conexões entre os neurônios. Mas isso também aumenta o consumo de energia e a geração de calor, tornando o cérebro maior e mais lento.

Portanto, cada vez que movemos uma peça do cérebro, levamos um xequemate. As leis da física parecem indicar que já chegamos ao auge da inteligência que podemos obter dessa forma. A menos que possamos aumentar subitamente o tamanho do crânio, ou a própria natureza dos neurônios, parece que chegamos ao limite máximo da inteligência. Se quisermos expandir a inteligência, a única maneira é tornar o cérebro mais eficiente (via remédios, genes e, possivelmente, equipamentos como o da EMT).

PENSAMENTOS FINAIS

Em suma, nas próximas décadas será possível usar uma combinação de terapia de genes, medicamentos e dispositivos magnéticos para expandir nossa inteligência. Há muitas vias de exploração revelando os segredos da inteligência e como ela poderá ser modificada ou melhorada. Mas o que acontecerá com a sociedade, se conseguirmos aprimorar a inteligência e dar um “salto cerebral”? Os especialistas em ética ponderaram seriamente sobre essa questão, já que a ciência básica vem crescendo com tanta rapidez. O medo maior é de que a sociedade possa se dividir, tendo apenas os ricos e poderosos com acesso a essa tecnologia que solidificaria ainda mais sua posição privilegiada na sociedade. E os pobres, sem acesso a um maior poder cerebral, teriam ainda mais dificuldade de ascensão social.

Essa é certamente uma preocupação válida, mas contraria a história da tecnologia. Muitas tecnologias do passado eram de fato destinadas inicialmente aos ricos e poderosos. Mas aos poucos, a produção em massa, a competição, os meios de transporte e os avanços tecnológicos fizeram cair os preços, e o cidadão comum passou a poder pagar por elas. (Por exemplo: achamos as refeições de hoje algo muito normal, mas nem o rei da Inglaterra podia tê-las um século atrás. A tecnologia tornou possível comprar em qualquer supermercado iguarias do mundo inteiro, que fariam inveja aos aristocratas da era vitoriana.) Portanto, se for possível expandir a inteligência, o preço da tecnologia também cairá gradualmente. A tecnologia nunca é monopólio dos ricos. Mais cedo ou mais tarde, a criatividade, o trabalho árduo e as próprias forças do mercado fazem baixar os preços.

Há também o medo de que a raça humana se divida entre os que querem desenvolver a inteligência e os que preferem mantê-la como está, resultando no pesadelo de termos uma classe de senhores bramânicos superinteligentes dominando as massas de menos dotados.

Mais uma vez, o medo de um salto na inteligência talvez seja exagerado. As pessoas em geral não têm o menor interesse em solucionar as complexas

equações tensoriais para um buraco negro. As pessoas em geral não têm nada a ganhar dominando a matemática das dimensões hiperespaciais nem a física da teoria quântica. As pessoas em geral, pelo contrário, podem achar essas atividades muito entediantes e inúteis. Assim, muitos não serão gênios da matemática, se tiverem essa oportunidade, porque não faz parte de sua personalidade, e eles não enxergam nenhuma vantagem nisso.

Tenhamos em mente que a sociedade já possui uma classe de matemáticos e físicos de boa formação, que ganham muito menos dinheiro do que qualquer empresário e exercem muito menos poder do que os políticos em geral. Ser superinteligente não garante sucesso financeiro. Na verdade, um superinteligente pode ser relegado aos degraus mais baixos de uma sociedade que valoriza tanto os atletas, estrelas de cinema, humoristas e apresentadores de televisão.

Ninguém jamais ficou rico estudando a relatividade.

E muito vai depender das características que serão ampliadas. Há outras formas de inteligência além da matemática. (Alguns defendem que inteligência inclui o talento artístico. Nesse caso, conseguimos imaginar o uso desse talento para levar uma vida confortável.)

Pais ansiosos de crianças em idade escolar podem querer incrementar o Q.I. dos filhos nos cursos preparatórios para os exames protocolares. Mas o Q.I., como vimos, não garante necessariamente sucesso na vida. Da mesma forma, muitas pessoas desejam melhorar a memória, mas, como vimos com os savants, ter uma memória fotográfica pode ser uma bênção e uma maldição. Nos dois casos, é improvável que o aprimoramento vá contribuir para uma divisão da sociedade.

A sociedade como um todo, porém, pode se beneficiar dessa tecnologia. Profissionais com maior inteligência estarão mais preparados para lidar com as mudanças frequentes do mercado de trabalho. A requalificação profissional será um gasto a menos para a sociedade. Além disso, o público será capaz de tomar decisões conscientes sobre questões tecnológicas importantes para o futuro (por exemplo, mudanças climáticas, energia nuclear, exploração espacial), porque saberá mais sobre esses assuntos complexos.

Tal tecnologia pode contribuir até para o campo das brincadeiras infantis. As crianças que hoje frequentam escolas particulares elitistas e têm professores particulares são mais preparadas para o mercado de trabalho porque têm mais chances de dominar questões difíceis. Mas se todo mundo tiver a inteligência melhorada, os abismos sociais serão ultrapassados. Assim, os objetivos alcançados pelo indivíduo estarão mais relacionados à sua inclinação, ambição, imaginação e recursos internos do que ao fato de ter nascido em uma família rica.

Além disso, a expansão da inteligência pode acelerar as inovações tecnológicas. Maior inteligência significa maior capacidade de simular o futuro, o

que é de valor inestimável para as descobertas científicas. Muitas vezes a ciência fica estagnada em certas áreas devido à falta de boas ideias para simular novos caminhos de pesquisa. A capacidade de simular diversos futuros possíveis aumenta bastante a probabilidade de descobertas científicas.

E as conquistas científicas, por sua vez, podem gerar novas indústrias, que enriquecerão toda a sociedade, criando novos mercados, empregos e oportunidades. A história está cheia de descobertas tecnológicas que criaram indústrias inteiramente novas, beneficiando não só alguns, mas toda a sociedade (por exemplo, o transistor e o laser, que são hoje a base da economia mundial).

Entretanto, na ficção científica sempre há um supermalfeitor que usa seu poder cerebral superior para arquitetar os crimes mais variados e frustrar o super-herói. Todo Super-Homem tem seu Lex Luthor, todo Homem-Aranha tem seu Duende Verde. Embora seja bem possível que uma mente criminososa use o aprimoramento cerebral para criar superarmas e planejar o crime do século, a polícia também pode ter a inteligência melhorada para deter o bandido. Assim, os supermalfeitores só serão perigosos se forem os únicos de posse da inteligência aumentada.

Até aqui, examinamos a possibilidade de expandir ou alterar nossas capacidades mentais por meio da telepatia, telecinesia, *upload* de memória e aprimoramento cerebral. Essas melhorias significam basicamente modificar e aumentar as aptidões mentais de nossa consciência. Isso supõe tacitamente que nossa consciência normal seja a única que existe, mas eu gostaria de investigar se há outras formas de consciência. Se houver, deve também haver outras maneiras de pensar, que levem a resultados e consequências totalmente diferentes. Em nossos pensamentos, há estados alterados de consciência, como sonhos, alucinações induzidas por drogas e doença mental. Há também a consciência inumana, a consciência dos robôs, e até a consciência dos alienígenas do espaço sideral. Precisamos descartar a noção chauvinista de que a consciência humana é a única existente. Há mais de um meio de criar um modelo do nosso mundo, e mais de um meio de simular seu futuro.

Os sonhos, por exemplo, são uma das formas mais antigas de consciência, e foram estudados já na antiguidade. No entanto, até recentemente houve muito pouco progresso em seu entendimento. Talvez os sonhos não sejam eventos bobos, aleatórios, costurados pelo cérebro adormecido, mas fenômenos que podem nos ajudar a penetrar nos significados da consciência. Os sonhos podem ser fundamentais para o entendimento dos estados alterados de consciência.

3. Trocadilho com o duplo sentido do termo *odd*: “As chances são boas, mas os bons são esquisitos.” (N. da T.)

LIVRO III ALTERAÇÕES DE CONSCIÊNCIA

O futuro pertence àqueles que acreditam na beleza dos seus sonhos.

– ELEANOR ROOSEVELT

7 EM SEUS SONHOS

Os sonhos podem determinar o destino.

Talvez o sonho mais famoso da antiguidade tenha ocorrido no ano 312 d.C., quando o imperador romano Constantino estava empenhado numa das maiores batalhas de sua vida. Ao encarar um exército inimigo duas vezes maior que o seu, ele percebeu que provavelmente morreria na batalha do dia seguinte. Mas naquela noite ele sonhou com um anjo trazendo a imagem de uma cruz e dizendo as proféticas palavras “Com este símbolo, vencerás”. Imediatamente, Constantino ordenou que todos os escudos fossem adornados com a imagem da cruz.

A história conta que ele saiu triunfante no dia seguinte, consolidando seu domínio à frente do Império Romano. Em vista dessa vitória, Constantino jurou pagar o débito de sangue que tinha para com a religião relativamente obscura chamada cristianismo, perseguida durante séculos pelos imperadores romanos, e cujos adeptos costumavam servir de alimento para os leões no Coliseu. Constantino assinou leis que iriam abrir o caminho para que o cristianismo se tornasse a religião oficial de um dos maiores impérios do mundo.

Durante milhares de anos, os sonhos intrigaram tanto reis e rainhas como mendigos e ladrões. Os antigos acreditavam que os sonhos eram presságios, e ao longo da história houve incontáveis tentativas de interpretá-los. A Bíblia relata, em Gênesis 41, a façanha de José, que interpretou corretamente o sonho do faraó do Egito, milhares de anos atrás. O faraó sonhou com sete vacas gordas, seguidas por sete vacas magras. Ficou tão aflito, que chamou escribas e místicos de todo o reino para encontrar o significado do sonho. Nenhum deles soube dar uma explicação satisfatória, até que José apareceu. Em sua interpretação do sonho, o Egito teria sete anos de boas colheitas e fartura, seguidos por sete anos de seca e fome. Assim, disse José, o Egito precisava começar a estocar os grãos e suprimentos, preparando-se para os anos de privação e desespero que viriam. Quando isso realmente aconteceu, José foi considerado um profeta.

Os sonhos sempre foram associados a profecia, mas em tempos mais recentes, são também conhecidos por estimular descobertas científicas. A ideia de que os neurotransmissores facilitam o movimento das informações por meio das sinapses, o que veio a ser a base da neurociência, surgiu num sonho do farmacologista Otto Loewi. Em 1865, August Kekulé sonhou com benzeno, uma molécula em que a ligação dos átomos de carbono forma uma cadeia que se fecha num círculo, como uma cobra mordendo o próprio rabo. Esse sonho desvendou a estrutura atômica da molécula de benzeno. Ele concluiu: “Vamos aprender a sonhar!”

Os sonhos também têm sido interpretados como uma janela para nossos verdadeiros pensamentos e intenções. Michel de Montaigne, o grande escritor e

ensaísta da Renascença, escreveu: “Acredito que os sonhos são a verdadeira interpretação de nossas inclinações, mas é preciso ter a arte de classificá-los e entendê-los.” Mais recentemente, Sigmund Freud propôs uma teoria para explicar a origem dos sonhos. Em sua reconhecida obra *A interpretação dos sonhos*, ele afirma que os sonhos são manifestações de desejos inconscientes, recalcados pela mente consciente, que se libertam durante o sono. Os sonhos não são apenas fragmentos aleatórios de imaginações inflamadas, mas revelam realmente segredos e verdades sobre o indivíduo. Na definição de Freud, “o sonho é a estrada real para o inconsciente”. Desde então, foram escritos inúmeros livros afirmando revelar o sentido oculto de qualquer imagem perturbadora sob a ótica da teoria freudiana.

Hollywood se aproveita de nossa eterna fascinação pelos sonhos. Uma cena predileta em muitos filmes é quando o herói tem um pesadelo aterrorizante e de repente acorda banhado em suor. No sucesso de bilheteria *A origem*, Leonardo DiCaprio é um ladrãozinho que rouba segredos íntimos do mais improvável de todos os lugares, os sonhos das pessoas. Munido de uma nova invenção, ele entra nos sonhos dos outros e os ilude para revelarem seus segredos financeiros. Corporações gastam milhões de dólares na proteção de segredos e patentes industriais. Bilionários escondem sua fortuna atrás de códigos complicados. E o trabalho dele é roubá-los. A trama se desenvolve rapidamente quando os personagens entram em sonhos nos quais uma pessoa adormece e sonha também. Assim os criminosos descem cada vez mais fundo nas múltiplas camadas do inconsciente.

Embora os sonhos sempre tenham nos causado fascínio e perplexidade, apenas na última década, ou pouco mais, os cientistas conseguiram penetrar nos mistérios oníricos. Na verdade, os cientistas realizam hoje algo que era considerado impossível: tiram fotos e gravam vídeos rudimentares de sonhos com aparelhos de IRM. Um dia você poderá ver um vídeo do que sonhou na noite anterior, e conhecer sua mente inconsciente. Com um treinamento adequado, você poderá fazer com que a consciência controle a natureza dos seus sonhos. E talvez, como o personagem de DiCaprio, uma tecnologia avançada permitirá que você entre no sonho alheio.

A NATUREZA DOS SONHOS

Por mais misteriosos que sejam, os sonhos não são um luxo supérfluo, ruminções inúteis de um cérebro ocioso. Na verdade, são essenciais para nossa sobrevivência. As varreduras cerebrais mostram que certos animais têm a atividade cerebral típica do sonho. Se forem privados do sonho, esses animais

podem morrer mais depressa do que se sucumbissem à fome, porque essa privação prejudica seriamente seu metabolismo. Infelizmente, a ciência não sabe dizer exatamente por que isso acontece.

Sonhar é um aspecto essencial do ciclo do sono. Passamos cerca de duas horas por noite sonhando, e cada sonho dura de cinco a vinte minutos. Durante a vida, passamos cerca de seis anos sonhando.

Os sonhos também têm as mesmas características em toda a raça humana. Pesquisando diferentes culturas, os cientistas encontraram temas comuns. O professor de psicologia Calvin Hall registrou 50 mil sonhos num período de 40 anos. E complementou o registro com mil relatos de estudantes universitários. Como ele esperava, constatou que a maioria das pessoas sonhava com as mesmas coisas, como experiências dos dias ou semanas anteriores. Os animais, porém, parecem sonhar de forma diferente de nós. No golfinho, por exemplo, apenas um hemisfério adormece de cada vez para evitar afogamento, pois sendo mamíferos e não peixes, eles precisam de ar para respirar. Então, se sonham, provavelmente é com um hemisfério de cada vez.

Como vimos, o cérebro não é um computador, e sim uma rede neural que se atualiza a cada tarefa aprendida. Os cientistas que trabalham com redes neurais observaram algo interessante. Frequentemente, os sistemas ficam saturados quando aprendem demais e, em vez de processar mais informação, entram num estado de “sonho”, em que memórias aleatórias às vezes aparecem e se fundem enquanto a rede neural tenta digerir o novo material. Então os sonhos podem equivaler a uma “faxina” em que o cérebro tenta organizar a memória de forma mais coerente. Se isso for verdade, é possível que todas as redes neurais, inclusive as de todos os organismos capazes de aprender, entrem num estado de sonho a fim de fazer uma arrumação das lembranças. Nesse caso, os sonhos provavelmente servem a um propósito. Alguns cientistas especulam que isso pode significar que robôs que aprendem com a experiência também podem sonhar.

Estudos neurológicos parecem sustentar essa conclusão. Alguns estudos mostraram que um período de sono entre a atividade e um teste pode melhorar a retenção de memória. As neuroimagens indicam que as áreas cerebrais ativadas durante o sono são as mesmas que as envolvidas no aprendizado de uma nova tarefa. Talvez sonhar seja útil para consolidar novas informações.

Os sonhos podem também trazer acontecimentos de poucas horas atrás, logo antes de dormir, mas geralmente trazem lembranças de alguns dias antes. Por exemplo: experimentos mostraram que se alguém usa óculos com lentes cor de rosa, alguns dias depois os sonhos ficam cor de rosa.

A varredura cerebral está desvendando um pouco dos mistérios dos sonhos. Normalmente, o EEG mostra que o cérebro emite ondas eletromagnéticas uniformes enquanto estamos acordados. No entanto, quando adormecemos gradualmente, os sinais do EEG começam a mudar de frequência. Depois, quando sonhamos, ondas de energia elétrica emanam do tronco encefálico e sobem para as áreas corticais, especialmente para o córtex visual. Isso confirma que as imagens visuais são um componente importante dos sonhos. Por fim, entramos no estado de sonho, e as ondas cerebrais são representadas pelo movimento rápido dos olhos (REM, *rapid eye movements*). (Dado que alguns mamíferos também têm a fase REM no sono, podemos deduzir que eles sonham.)

Enquanto as áreas visuais do cérebro estão ativas, áreas ligadas a odor, sabor e tato são desligadas. Quase todas as imagens e sensações processadas pelo corpo são autogeradas, tendo origem nas vibrações eletromagnéticas do tronco encefálico, e não em estímulos externos. O corpo fica muito isolado do mundo externo e, quando sonhamos, ficamos mais ou menos paralisados. Talvez essa paralisia é que nos impeça de agir fisicamente conforme o sonho, o que poderia ser desastroso. Cerca de 6% das pessoas sofrem do distúrbio de “paralisia do sono”, em que despertam do sonho ainda paralisadas. Muitas vezes, essas pessoas acordam assustadas, achando que há criaturas prendendo seu peito, braços, pernas. Quadros da era vitoriana retratam mulheres acordando com um duende horroroso sentado em cima delas, encarando-as. Alguns psicólogos acreditam que a paralisia do sono pode explicar a origem da síndrome da abdução alienígena.

O hipocampo fica ativo quando sonhamos, o que sugere que os sonhos utilizam nosso estoque de lembranças. A amígdala e o cíngulo anterior também ficam ativos, o que significa que os sonhos podem trazer grandes emoções, inclusive o medo.

Ainda mais reveladoras, porém, são as áreas cerebrais que se desligam, como o córtex pré-frontal dorsolateral (que é o centro de comando do cérebro), o córtex orbitofrontal (que age como censor, ou verificador de fatos) e a região temporal-parietal (que processa os sinais sensorio-motores e o senso espacial).

Quando o córtex pré-frontal dorsolateral está desligado, não podemos contar com o centro de planejamento racional. Então vagamos a esmo nos sonhos, com o centro visual fornecendo imagens sem o controle racional. Como o córtex orbitofrontal, que verifica os fatos, também está inativo, os sonhos fluem livremente, sem as restrições das leis da física ou do senso comum. E o desligamento do lobo temporal-parietal, que contribui para coordenar o senso de onde estamos por meio de sinais enviados pelo ouvido interno e pelos olhos, pode explicar as experiências fora do corpo quando sonhamos.

Já enfatizamos que a consciência humana representa principalmente o

cérebro, sempre criando modelos do mundo externo e os simulando no futuro. Nesse caso, os sonhos representam um caminho alternativo de simulação do futuro, em que as leis da natureza e as interações sociais são suspensas temporariamente.

COMO SONHAMOS?

Mas as explicações acima não respondem à seguinte questão: Como os sonhos são gerados? Uma das maiores autoridades no assunto é o dr. Allan Hobson, psiquiatra da Harvard Medical School, que dedicou décadas de sua vida a desvendar os mistérios dos sonhos. Ele afirma que os sonhos, especialmente na fase REM, podem ser estudados em nível neurológico, e que eles surgem quando o cérebro tenta dar sentido aos sinais aleatórios que emanam do tronco cerebral.

Quando o entrevistei, dr. Hobson disse que, após décadas de catalogação de sonhos, ele encontrou cinco características básicas:

1. Emoções intensas – isso se deve à ativação da amígdala, provocando emoções como o medo.
2. Conteúdo ilógico – os sonhos mudam imediatamente de uma cena para outra, desafiando a lógica.
3. Supostas impressões sensoriais – os sonhos nos dão sensações falsas, que são geradas internamente.
4. Aceitação acrítica de eventos do sonho – aceitamos, sem crítica, a natureza ilógica do sonho.
5. Dificuldade de lembrar – os sonhos são logo esquecidos, minutos após o despertar.

O dr. Hobson, junto com o dr. Robert McCarley, fez história ao propor o primeiro desafio sério à teoria freudiana dos sonhos, denominado “hipótese da ativação-síntese”. Em 1977, eles propuseram a ideia de que os sonhos se originam de disparos aleatórios de neurônios do tronco cerebral, subindo para o córtex que, por sua vez, tenta dar sentido a esses sinais aleatórios.

A chave dos sonhos está em núcleos localizados no tronco cerebral, a parte mais antiga do cérebro, que segregam substâncias químicas especiais, chamadas adrenérgicas, que nos mantêm alertas. Quando vamos dormir, o tronco cerebral ativa outro sistema, o colinérgico, que emite substâncias químicas que nos

colocam no estado de sonho.

Quando sonhamos, os neurônios colinérgicos do tronco cerebral entram em atividade, emitindo pulsos erráticos de energia elétrica chamados ondas PGO (ponto-genículo-occipital). Essas ondas sobem pelo tronco cerebral e estimulam o córtex visual a criar os sonhos. As células do córtex visual entram em ressonância, vibrando centenas de vezes por segundo de maneira irregular, o que talvez seja responsável pela natureza incoerente dos sonhos.

Esse sistema também emite substâncias químicas que desativam as partes do cérebro responsáveis pela razão e pela lógica. A falta de verificação pelos córtices pré-frontal e orbitofrontal, juntamente com a extrema sensibilidade do cérebro a pensamentos errantes podem explicar a natureza estranha, errática, dos sonhos.

Estudos mostram que é possível entrar no estado colinérgico sem dormir. O dr. Edgard Garcia-Rill, da Universidade do Arkansas, afirma que meditação, preocupação, ou ficar num tanque de isolamento sensorial podem induzir ao estado colinérgico. Pilotos e motoristas, que passam muitas horas diante da monotonia do para-brisa, sem imagens, também podem entrar nesse estado. Em sua pesquisa, ele descobriu que esquizofrênicos têm uma quantidade incomum de neurônios colinérgicos no tronco cerebral, o que poderia explicar algumas das alucinações.

Para dar mais consistência a seu estudo, o dr. Allan Hobson mandou suas cobaias dormirem com uma touca especial que registra automaticamente os dados durante o sonho. Um sensor conectado à touca registra os movimentos da cabeça (porque os movimentos da cabeça geralmente ocorrem quando o sonho termina). Outro sensor mede os movimentos das pálpebras (porque o REM faz as pálpebras se mexerem). Quando os sujeitos acordavam, relatavam imediatamente o que tinham sonhado, e as informações da touca eram colocadas no computador.

Dessa maneira, o dr. Hobson acumulou uma grande quantidade de informação sobre os sonhos. Mas qual é o significado dos sonhos?, perguntei. Ele descarta o que chama de “interpretação mística do tipo biscoito da sorte chinês”. E não vê nos sonhos nenhuma mensagem oculta vinda do cosmos.

Ele acredita que, após as ondas PGO subirem pelo tronco cerebral para as áreas corticais, o córtex tenta dar sentido aos sinais erráticos, e inventa uma narrativa a partir deles: o sonho.

A FOTOGRAFIA DE UM SONHO

No passado, quase todos os cientistas evitavam estudar os sonhos, porque são

muito subjetivos e têm uma longa história associada ao misticismo e ao psiquismo. Mas com a IRM, os sonhos estão revelando seus segredos. De fato, como os centros cerebrais que controlam o sonho são quase idênticos aos que controlam a visão, é possível fotografar o sonho. Esse trabalho pioneiro está sendo feito em Kioto, no Japão, por cientistas do ATR Computational and Neuroscience Laboratories.

Os sujeitos são colocados no aparelho de IRM, onde lhes são apresentadas 400 imagens em preto e branco, cada uma consistindo em um conjunto de pontos num quadro de 10x10 pixels. As imagens são projetadas uma a uma e a IRM registra a resposta do cérebro a cada coleção de pixels. Assim como outros grupos que trabalham nessa área de interface cérebro-máquina, os cientistas criam uma enciclopédia, onde cada imagem corresponde a um único padrão de IRM. Dessa maneira, os cientistas podem trabalhar no sentido inverso para reconstruir corretamente as imagens autogeradas na varredura por IRM enquanto o sujeito estava sonhando.

Yukiyasu Kamitani, cientista chefe do ATR, disse: “Essa tecnologia pode também ser aplicada a outros sentidos além da visão. No futuro, pode até ser possível ler sentimentos e estados emocionais complicados.” Assim, qualquer estado mental pode ser transformado em imagem, inclusive os sonhos, desde que se faça um glossário dos estados mentais e imagens de IRM correspondentes.

Os cientistas de Kioto se concentraram em analisar fotografias geradas pela mente. No capítulo 3, vimos uma abordagem semelhante, introduzida por Jack Gallant, em que, com ajuda de uma fórmula complexa, os vóxeis da varredura por IRM em 3D podem ser usados para reconstruir a imagem vista pelo olho. Um processo similar permitiu ao dr. Gallant e sua equipe criar um vídeo rudimentar de um sonho. Quando visitei seu laboratório em Berkeley, conversei com um membro da equipe de pós-doutorado, dr. Shinji Nishimoto, que me permitiu assistir ao vídeo de um sonho, um dos primeiros a serem feitos. Vi uma série de rostos tremulando na tela do computador, significando que a cobaia (nesse caso, o próprio dr. Nishimoto) estava sonhando com pessoas, e não com animais ou objetos. É extraordinário. Infelizmente, a tecnologia ainda não é suficientemente boa para se ver com exatidão os traços das pessoas que aparecem no sonho, portanto, o próximo passo será aumentar o número de pixels, para que as imagens fiquem mais complexas. Outro avanço será reproduzir imagens em cores em vez de preto e branco.

Fiz ao dr. Nishimoto a pergunta crucial: Como você sabe que o vídeo é fiel? Como você sabe que o aparelho não está inventando coisas? Ele ficou um pouco encabulado ao responder que esse era um ponto fraco da pesquisa. Normalmente, o registro do sonho só é possível durante alguns minutos após o despertar. Depois disso, muitos sonhos se perdem nas névoas da consciência, e não é fácil verificar os resultados.

O dr. Gallant me disse que essa pesquisa sobre gravação de sonhos estava em desenvolvimento e que ainda não tinha condições de ser publicada. Resta um longo caminho a percorrer antes que possamos assistir a um vídeo do sonho da noite anterior.

SONHOS LÚCIDOS

Os cientistas estão investigando também uma forma de sonho que já se pensou ser um mito: o sonho lúcido, ou sonhar em estado consciente. Parece uma contradição de termos, mas foi verificado em varreduras cerebrais. No sonho lúcido, a pessoa sabe que está sonhando e pode, conscientemente, controlar a direção do sonho. Embora a ciência só tenha começado a fazer experimentos com sonhos lúcidos recentemente, há referências de séculos anteriores a esse fenômeno. No budismo, por exemplo, existem livros que citam sonhadores lúcidos e até ensinam a se tornar um deles. No decorrer do tempo, várias pessoas na Europa escreveram relatos detalhados de sonhos lúcidos.

Varreduras cerebrais de sonhadores lúcidos mostram que o fenômeno é real. Durante a fase REM, o córtex pré-frontal dorsolateral dessas pessoas, geralmente adormecido quando alguém normal sonha, continua ativo, indicando certa consciência enquanto ele sonha. De fato, quanto mais lúcido é o sonho, mais ativo fica o córtex pré-frontal dorsolateral. Como esse córtex representa a parte consciente do cérebro, o sonhador deve ter consciência de que está sonhando.

O dr. Hobson me disse que qualquer pessoa pode aprender a ter sonhos lúcidos com a utilização de certas técnicas. Especialmente quem tem sonhos lúcidos deve manter um caderno para anotar o que sonha. Antes de dormir, devem se lembrar de que irão “acordar” no meio do sonho, sabendo que estão se movimentando num mundo de sonhos. É importante ter isso em mente antes de pôr a cabeça no travesseiro. Como o corpo fica bem paralisado na fase REM, é difícil enviar ao mundo externo um sinal de que se começou a sonhar, mas o dr. Stephen LaBerge, da Universidade de Stanford, estudou sonhadores lúcidos (inclusive ele próprio) que conseguem sinalizar para o mundo externo que estão sonhando.

Em 2011, cientistas usaram pela primeira vez sensores de IRM e EEG para medir o conteúdo de sonhos, e até fizeram contato com uma pessoa sonhando. No Max Planck Institute, em Munique e Leipzig, cientistas tiveram ajuda de sonhadores lúcidos equipados com sensores de EEG para determinar o momento em que entravam na fase REM. E então eram colocados no aparelho de IRM. Antes de adormecer, os sonhadores concordaram em fazer uma série de movimentos de olhos e respiração, como um código Morse, enquanto sonhavam.

Quando começassem a sonhar, deveriam cerrar o punho direito e depois o esquerdo durante dez segundos. Era o sinal de que estavam sonhando.

Os cientistas observaram que quando as cobaias começavam a sonhar, o córtex sensorio-motor (responsável pelo controle de ações motoras, como apertar os punhos) era ativado. A IRM captava que os punhos estavam sendo cerrados, e qual punho era cerrado primeiro. Depois, usando outro sensor (um espectrômetro quase infravermelho), confirmaram que havia maior atividade na região cerebral que controla o planejamento de movimentos.

“Nossos sonhos, portanto, não são um ‘filme do sono’, ao qual apenas assistimos passivamente; eles exigem atividade em regiões cerebrais relevantes para o seu conteúdo”, diz Michael Czisch, líder de uma equipe de pesquisa no Max Planck Institute.

ENTRANDO NUM SONHO

Se podemos nos comunicar com uma pessoa que está sonhando, então é também possível interferirmos no sonho de alguém? Sim, provavelmente.

Como vimos, os cientistas já deram os primeiros passos na gravação de imagens de sonhos, e nos próximos anos será possível obter fotos e vídeos com uma precisão muito maior. Como os cientistas já conseguem estabelecer um canal de comunicação entre o mundo real e um sonhador lúcido no mundo da fantasia, em princípio, os cientistas serão capazes de modificar o curso do sonho. Digamos que cientistas estejam assistindo ao vídeo de um sonho usando um aparelho de IRM enquanto o sonho se desenrola em tempo real. Enquanto a pessoa vaga pelo mundo dos sonhos, os cientistas podem dizer para onde ela está indo e orientá-la sobre uma direção diferente.

Assim, num futuro próximo, talvez seja possível assistir ao vídeo de um sonho e influenciar seu curso. Mas no filme *A origem*, Leonardo DiCaprio vai muito além. Ele é capaz não só de assistir ao sonho de outras pessoas, mas de entrar no sonho delas. Será possível?

Já vimos que ficamos paralisados quando sonhamos, portanto, não colocamos em prática as fantasias do sonho, o que seria desastroso. Contudo, os sonâmbulos geralmente andam dormindo com os olhos abertos (embora o olhar pareça embaçado). Desse modo, eles vivem num mundo híbrido, meio real, meio sonho. Há muitos casos documentados de pessoas andando em volta da casa, dirigindo carro, cortando lenha, e até cometendo homicídio nesse estado de sonho, onde os mundos da realidade e da fantasia se misturam. Assim, é possível que as imagens físicas que o olho enxerga possam de fato interagir livremente com as imagens fictícias que o cérebro fabrica durante um sonho.

A maneira de entrar no sonho alheio, portanto, deve ser colocar na cobaia lentes de contato que projetem as imagens diretamente em sua retina. Protótipos de lentes de contato conectadas à internet já estão sendo desenvolvidos na Universidade de Washington em Seattle. Assim, se o observador quiser entrar nesse sonho, ele precisa ficar num estúdio sendo filmado por uma câmera de vídeo. Sua imagem será projetada nas lentes de contato do sonhador, criando uma imagem composta (a imagem do observador superposta à cena imaginária que o cérebro está produzindo).

O observador pode então ver realmente o mundo sonhado onde ele agora está presente, porque ele também está usando lentes de contato conectadas à internet. A imagem em IRM do sonho da cobaia é decifrada pelo computador e enviada diretamente para as lentes de contato do observador.

Assim, o observador pode mudar a direção do sonho em que ele entrou. Andando pelo estúdio vazio, o observador poderá ver o sonho se desenrolar em suas lentes de contato e interagir com os objetos e pessoas que aparecem nele. Será uma experiência incrível, pois o cenário muda de repente, imagens aparecem e desaparecem sem motivo algum, e as leis da física estão suspensas. Vale tudo.

Num futuro ainda mais distante, talvez seja possível entrar no sonho de alguém conectando dois cérebros em estado de sono. Cada cérebro teria que ser conectado a um aparelho de IRM ligado a um computador central, que iria fundir os dois sonhos num só. Primeiro, o computador teria que transformar a varredura cerebral de cada um em imagem de vídeo. Depois, o sonho de um deles seria enviado para as áreas sensoriais do cérebro do outro, de modo que o sonho do segundo se fundisse com o do primeiro. Entretanto, a tecnologia de gravação de vídeo e de interpretação de sonhos terá que avançar muito até isso se tornar uma possibilidade real.

Mas isso levanta outra questão: Se é possível alterar o curso do sonho de alguém, então é possível controlar não só o sonho, mas também a mente de outra pessoa? Durante a Guerra Fria, isso foi um tema de extrema gravidade, com os Estados Unidos e a União Soviética empenhados num jogo mortal, tentando usar técnicas psicológicas para controlar a vontade do outro.

A mente é só o que o cérebro faz.

– MARVIN MINSKY

8 A MENTE PODE SER CONTROLADA?

Um touro enfurecido adentra a arena vazia em Córdoba, na Espanha. Durante gerações, esses ferozes animais têm sido criados para maximizar seus instintos de matar. Então um professor de Yale entra calmamente na arena. Em vez de usar um paletó de tweed, ele está vestido como um toureiro corajoso, com uma vistosa jaqueta justa dourada. Ele abana a capa vermelha e provoca o touro. O professor não parece aterrorizado, mas sim tranquilo, confiante, até desligado. Para um observador, parece que o professor ficou maluco e quer se suicidar.

Enfurecido, o touro se depara com o professor e subitamente ataca, mirando-o com seus chifres mortais. O professor não sai correndo de medo. Ele tem uma caixinha na mão. Diante das câmeras, aperta um botão e o touro para no mesmo instante. O professor estava tão confiante que arriscou a vida para provar que dominava a arte de controlar a mente de um touro raivoso.

O professor de Yale era o dr. José Delgado, que estava anos à frente de seu tempo. Nos anos 1960, foi pioneiro em uma série de experimentos notáveis, porém perturbadores, com animais em cujo cérebro ele colocava eletrodos para controlar seus movimentos. A fim de fazer o touro parar, ele inseriu eletrodos no estriado do gânglio basal, na base do cérebro do animal, que é ligado à coordenação motora.

O professor fez uma série de experimentos com macacos, tentando reorganizar sua hierarquia social ao apertar um botão. Depois de implantar eletrodos no núcleo caudado (uma região associada ao controle motor) do macho alfa de um grupo, Delgado conseguiu reduzir a agressividade do líder. Os machos delta, sem a ameaça de retaliação, começaram a se impor, tomando o território e os privilégios normalmente exclusivos do alfa. Enquanto isso, o macho alfa parecia ter perdido o interesse em defender seu território.

Depois o dr. Delgado apertou outro botão e o macho alfa voltou ao normal instantaneamente, reassumindo o comportamento agressivo e o lugar de rei do pedaço. Os machos delta fugiram com medo.

Delgado foi o primeiro na história a demonstrar que era possível controlar a mente de animais dessa maneira. Tornou-se um titereiro, manipulando os cordões de marionetes vivas.

Como era de se esperar, a comunidade científica viu com desconforto o trabalho de Delgado. Para piorar a situação, ele escreveu um livro, em 1969, com o provocativo título *Physical Control of the Mind: Toward a Psychocivilized Society*, que levantou uma questão incômoda: se cientistas como o dr. Delgado estão manipulando as cordinhas, quem vai controlar os titereiros?

O trabalho de Delgado põe em foco os enormes perigos e promessas da tecnologia. Nas mãos de um ditador inescrupuloso, essa tecnologia pode ser usada para manipular e controlar seus infelizes subordinados. Mas ela também

pode ser usada para libertar milhões de pessoas aprisionadas em doenças mentais, aterrorizadas por alucinações ou consumidas por ansiedades.

Anos mais tarde, um jornalista perguntou ao dr. Delgado por que ele tinha iniciado aqueles experimentos controversos. Ele respondeu que desejava corrigir os horrendos abusos sofridos pelos doentes mentais. Eles eram submetidos a lobotomias radicais, em que o córtex pré-frontal é dilacerado por uma faca pontuda, como um furador de gelo, e enfiada a marteladas no cérebro logo acima do olho. Em geral, os resultados eram trágicos, e alguns dos horrores foram expostos no romance de Ken Kesey, *Um estranho no ninho*, que virou filme estrelado por Jack Nicholson. Alguns pacientes ficavam calmos e relaxados, mas muitos outros se tornavam zumbis, letárgicos, indiferentes à dor e aos sentimentos, emocionalmente vazios. Essa prática era tão difundida que, em 1949, Antonio Moniz ganhou o Prêmio Nobel por aperfeiçoar a lobotomia. Em 1950, a União Soviética proibiu essa tecnologia, declarando-a “contrária aos princípios de humanidade”, e a acusando de transformar “um insano num idiota”. No total, estima-se que foram realizadas 40 mil lobotomias somente nos Estados Unidos, durante duas décadas.

O CONTROLE DA MENTE E A GUERRA FRIA

Outro motivo para a falta de receptividade ao trabalho do dr. Delgado foi o clima político da época. A Guerra Fria estava no auge, com lembranças dolorosas de soldados norte-americanos prisioneiros exibidos diante das câmeras durante a Guerra da Coreia. Esses soldados, com o olhar vazio, admitiam ser espiões em missões secretas, confessavam horrendos crimes de guerra e denunciavam o imperialismo dos Estados Unidos.

Para dar sentido a essas cenas, a imprensa usava o termo “lavagem cerebral”, com a ideia de que os comunistas tinham desenvolvido drogas e técnicas secretas para transformar soldados em zumbis manipuláveis. Nesse clima político carregado, Frank Sinatra estrelou, em 1962, o filme de suspense sobre a Guerra Fria intitulado *Sob o domínio do mal*, em que ele tenta expor um agente secreto comunista infiltrado cuja missão é assassinar o presidente dos Estados Unidos. Mas a trama dá uma virada. O verdadeiro assassino é um respeitado herói de guerra norte-americano que foi capturado e submetido a lavagem cerebral pelos comunistas. Nascido numa família muito bem relacionada, o agente secreto está acima de qualquer suspeita, e é quase impossível detê-lo. *Sob o domínio do mal* reflete o medo dos norte-americanos naquela época.

Muitos desses medos foram retratados no profético romance de 1931 de Aldous Huxley, *Admirável mundo novo*. Nessa distopia há grandes fábricas de

bebês de proveta que produzem clones. Privando os fetos de oxigênio, seletivamente, conseguem produzir crianças com diversos graus de lesão cerebral. No topo estão os alfas, que não sofrem de lesões no cérebro e foram criados para comandar a sociedade. No nível mais baixo estão os ipsilons, que sofreram danos cerebrais consideráveis e são usados como operários obedientes e descartáveis. Nos níveis intermediários estão outras categorias de trabalhadores e os burocratas. A elite controla a população com drogas que alteram a mente, com amor livre e lavagem cerebral. Desse modo, a paz, a tranquilidade e a harmonia são mantidas, mas a história traz uma questão perturbadora, que tem ressonâncias até os dias de hoje: quanto da nossa liberdade e dos princípios básicos de humanidade queremos sacrificar em nome da paz e da ordem social?

EXPERIMENTOS DA CIA EM CONTROLE DA MENTE

A histeria da Guerra Fria chegou aos mais altos escalões da CIA. Considerando os soviéticos muito mais avançados no domínio da lavagem cerebral e de métodos científicos não ortodoxos, a CIA se dedicou a vários projetos sigilosos, como o MKULTRA, que teve início em 1953, para explorar ideias grotescas e marginais. (Em 1973, quando o escândalo de Watergate espalhou pânico no governo, o diretor da CIA, Richard Helms, cancelou o MKULTRA e mandou destruir imediatamente todos os documentos ligados ao projeto. No entanto, 20 mil documentos escondidos sobreviveram e foram divulgados em 1977, sob o Freedom of Information Act, revelando todo o alcance daquela ação.)

Hoje sabe-se que, entre 1953 e 1973, o MKULTRA financiou 80 instituições, inclusive 44 universidades e faculdades, vários hospitais, laboratórios farmacêuticos e presídios, e foram feitos experimentos em pessoas sem a autorização delas, em 150 operações secretas. A certa altura, 6% do orçamento de toda a CIA era destinado ao MKULTRA.

Alguns desses projetos de controle da mente eram:

- o desenvolvimento de um “soro da verdade” que faria os prisioneiros confessarem seus segredos;
- apagamento de memória, num projeto da Marinha dos Estados Unidos chamado “Subproject 54”;
- uso de hipnose e de várias drogas, principalmente LSD, para controle de comportamento;
- pesquisa em uso de drogas de controle da mente contra líderes

estrangeiros, por exemplo, Fidel Castro;

- aperfeiçoamento de vários métodos de interrogatório de prisioneiros;
- desenvolvimento de uma droga fatal de efeito rápido e sem vestígios;
- alteração da personalidade por meio de drogas, para deixar as pessoas mais maleáveis.

Embora muitos cientistas questionassem a validade desses estudos, outros prosseguiram nessa linha. Foram recrutados especialistas de diversas áreas, abrangendo físicos, médiuns e cientistas da computação, para investigar uma variedade de projetos não ortodoxos, desde experimentos com drogas que alteram a mente, como o LSD, até a missão dada a médiuns de localizar submarinos soviéticos patrulhando as profundezas do mar, entre outros. Num triste incidente, deram secretamente LSD a um cientista do Exército dos Estados Unidos. Segundo alguns relatórios, ele ficou tão desorientado e violento que cometeu suicídio, se atirando de uma janela.

A maioria desses experimentos era justificada pela ideia de que os soviéticos estavam mais avançados que os norte-americanos em termos de controle da mente. O Senado dos Estados Unidos recebeu um relatório secreto dizendo que os soviéticos estavam fazendo experimentos com radiação de micro-ondas diretamente no cérebro de cobaias humanas. Em vez de denunciar essa ação, os Estados Unidos viram ali um “grande potencial de desenvolvimento de um sistema para desorientar ou romper os padrões de comportamento de militares e diplomatas”. O Exército dos norte-americanos chegou a afirmar que seria possível incutir palavras e até discursos inteiros na mente do inimigo: “Um conceito de isca e armadilha (...) é criar, por um meio remoto, ruídos na cabeça dos indivíduos, expondo-os a pulsos de micro-ondas de baixa potência. (...) É possível criar um discurso inteligível escolhendo-se adequadamente as características dos pulsos. (...) Assim pode ser possível ‘falar’ com determinados adversários de um modo que os desconcerte ao máximo”, dizia o relatório.

Infelizmente, nenhum desses experimentos foi verificado por especialistas, de modo que milhões de dólares dos contribuintes foram gastos em projetos como esse, que muito provavelmente violam as leis da física, pois o cérebro humano não pode receber radiação de micro-ondas e, principalmente, não tem capacidade de decodificar mensagens desse tipo. O dr. Steve Rose, biólogo da Open University, classificou essa trama esdrúxula de uma “impossibilidade neurocientífica”.

Apesar dos milhões de dólares gastos nesses “projetos escusos”, deles não resultou qualquer descoberta científica confiável. Na verdade, o uso de drogas para alterar a mente criou desorientação, e até pânico, nas cobaias, mas o

Pentágono não atingiu a principal meta, que era controlar a mente consciente de outra pessoa.

Segundo o psicólogo Robert Jay Lifton, a lavagem cerebral realizada pelos comunistas tinha pouco efeito em longo prazo. A maioria dos soldados americanos que denunciaram os Estados Unidos durante a Guerra da Coreia recuperou a personalidade anterior pouco após serem libertados. Além disso, estudos com pessoas que sofreram lavagem cerebral em certas seitas mostram que elas retomam a personalidade original após se desligarem da seita. Tudo indica que, no final das contas, a personalidade básica não é afetada pela lavagem cerebral.

Certamente, os militares não foram os primeiros a fazer experimentos de controle da mente. No passado, feiticeiros e videntes davam poções para fazer com que os capturados falassem ou se voltassem contra seus líderes. Um dos métodos mais antigos de controle da mente foi o hipnotismo.

VOCÊ ESTÁ FICANDO COM SONO, MUITO SONO...

Lembro-me de, quando criança, assistir a programas de televisão sobre hipnose. Num dos programas, uma pessoa em estado hipnótico recebeu a sugestão de que seria uma galinha ao despertar do transe. O público ficou abismado ao ver a pessoa batendo os braços e cacarejando no palco. Por mais dramática que tenha sido essa demonstração, era simplesmente um exemplo de “hipnose teatral”. Livros escritos por mágicos profissionais e apresentadores de programas de televisão revelam que eles utilizam atores infiltrados na plateia, usam a força da sugestão, e até a disposição da vítima para fazer parte do estratagema.

Uma vez apresentei um documentário de TV da BBC/Discovery chamado *Time*, e surgiu o tema das lembranças perdidas há muito tempo. É possível evocar lembranças tão distantes através da hipnose? E se for, é possível alguém impor sua vontade sobre a outra pessoa? Para testar essas ideias, deixei-me hipnotizar para um programa de TV.

A BBC contratou um hipnotizador profissional muito competente. Ele pediu que eu me deitasse, numa sala escura e silenciosa, e falou comigo em tom suave, lentamente, fazendo-me relaxar. Em seguida, me disse para pensar em algo do passado, algum lugar ou incidente que ainda se destacava mesmo depois de tantos anos. Disse-me para reentrar naquele lugar, vivenciando tudo o que vi ali, os sons, os cheiros. Para minha surpresa, comecei a ver lugares e pessoas de que eu me esquecera havia décadas. Era como assistir a um filme desfocado que ia ficando nítido aos poucos. Mas então a recordação terminou. Em certo ponto, não consegui me lembrar de mais nada. Havia claramente um limite ao que a

hipnose podia fazer.

Varreduras por EEG e IRM mostram que durante a hipnose o sujeito tem um mínimo de estimulação sensorial externa nos córtices sensoriais. Assim, a hipnose facilita o acesso a algumas lembranças já enterradas, mas certamente não muda a personalidade, os objetivos nem os desejos de alguém. Um documento secreto do Pentágono, de 1966, corrobora isso, explicando que o hipnotismo não é uma arma de guerra confiável. “É provavelmente significativo que, na longa história da hipnose, cuja aplicação potencial ao serviço de inteligência sempre foi conhecida, não existam relatos confiáveis de seu uso efetivo por um serviço de inteligência”, diz o documento.

Vale notar também que as varreduras cerebrais mostram que a hipnose não é um novo estado de consciência, como o sonho na fase REM. Se definirmos a consciência humana como um processo de construção contínua de modelos do mundo externo e então simularmos como esses modelos evoluem no futuro para atingir um objetivo, veremos que a hipnose não altera esse processo básico. A hipnose pode acentuar certos aspectos da consciência e ajudar a recuperar lembranças, mas não pode fazer alguém cacarejar como galinha sem o consentimento da própria pessoa.

DROGAS QUE ALTERAM A CONSCIÊNCIA E SOROS DA VERDADE

Uma das metas do MKULTRA era a criação de um soro da verdade que levasse espões e prisioneiros a revelar segredos. Embora o MKULTRA tenha sido extinto em 1973, os manuais de interrogatório da CIA e do Exército dos EUA tornados públicos pelo Pentágono em 1996 ainda recomendavam o uso de soros da verdade (apesar de o Supremo Tribunal dos Estados Unidos determinar que uma confissão obtida por esse meio é uma “coerção inconstitucional” e portanto inadmissível num julgamento).

Quem assiste a filmes de Hollywood sabe que o pentotal sódico é o soro da verdade predileto dos espões (como nos filmes *True Lies*, com Arnold Schwarzenegger, e *Entrando numa fria maior ainda*, com Robert De Niro). A droga faz parte de uma classe maior de barbitúricos, sedativos e hipnóticos capazes de ultrapassar a barreira sangue-cérebro que impede substâncias químicas nocivas na corrente sanguínea de penetrarem no cérebro.

É por isso que drogas que alteram a mente, como o álcool, nos afetam tanto, pois atravessam essa barreira. O pentotal sódico reduz a atividade do córtex pré-frontal e a pessoa fica mais relaxada, falante e desinibida. No entanto, isso não significa que fale a verdade. Pelo contrário, as pessoas sob a influência da substância, tal como as que beberam demais, são perfeitamente capazes de

mentir. Os “segredos” que saem da boca de alguém sob o efeito dessa droga podem ser grandes mentiras, e por isso até a CIA desistiu de usá-la.

Contudo, isso não exclui a possibilidade de algum dia descobrirem uma droga espetacular que possa alterar nossa consciência básica. Essa droga poderia alterar as sinapses entre as fibras nervosas, agindo sobre os neurotransmissores que operam nessa área, como a dopamina, a serotonina ou a acetilcolina. Se pensarmos nas sinapses como uma série de cabines de pedágio numa rodovia, certas drogas (estimulantes como a cocaína, por exemplo) poderiam abrir a cancela do pedágio, deixando passar livremente as mensagens. A súbita aceleração mental que os drogados sentem ocorre porque todas as cancelas são abertas ao mesmo tempo, provocando uma avalanche de sinais. Mas quando todas as sinapses disparam em uníssono, só podem ser ativadas de novo horas depois. É como se todas as barreiras de pedágio se fechassem, o que causa a súbita depressão que se segue à aceleração. O desejo do corpo de vivenciar novamente a aceleração é que causa o vício.

COMO AS DROGAS ALTERAM A MENTE

A base bioquímica das drogas que alteram a mente não era conhecida quando a CIA fez os experimentos em cobaias desavisadas, mas, desde então, a base molecular do vício em drogas tem sido muito pesquisada. Estudos com animais mostraram a força do vício em drogas: ratos, camundongos ou primatas, tendo oportunidade, usam drogas como cocaína, heroína e anfetaminas até caírem de exaustão ou morrerem em consequência do uso.

Para avaliar a extensão desse problema, considere que, em 2007, 13 milhões de norte-americanos, a partir de 12 anos (5% da população de adolescentes e adultos do país), experimentou ou se viciou em metanfetaminas. O vício em drogas não só destrói uma vida, mas destrói sistematicamente o cérebro. Varreduras por IRM de viciados em metanfetamina mostram uma redução de 11% do tamanho do sistema límbico, que processa emoções, e 8% de perda de tecido no hipocampo, que é o portal da memória. A IRM mostra que o dano é de certa forma comparável ao encontrado em pacientes de Alzheimer. Mas, apesar da destruição que a metanfetamina causa no cérebro, os viciados persistem no vício porque produz uma sensação até 12 vezes maior do que uma refeição deliciosa ou sexo.

Basicamente, o “barato” da droga deve-se ao fato de a substância capturar o próprio sistema de prazer/satisfação do cérebro, localizado no sistema límbico. Esse circuito de prazer/satisfação é muito primitivo, datando de milhões de anos de história evolucionária, mas ainda é extremamente importante para a

sobrevivência humana, porque premia o comportamento benéfico e pune as ações prejudiciais. Quando esse circuito é assumido pelas drogas, o resultado pode ser devastador. As drogas furam a barreira sangue-cérebro e causam uma superprodução de neurotransmissores como a dopamina, que inunda o núcleo accumbens, um pequeno centro de prazer perto da amígdala. A dopamina, por sua vez, é produzida por certas células cerebrais na área tegmentar ventral, chamadas células da ATV.

Todas as drogas funcionam basicamente da mesma maneira, bloqueando o circuito ATV-núcleo accumbens, que controla o fluxo de dopamina e outros neurotransmissores para o centro de prazer. As drogas diferem somente no modo pelo qual esse processo ocorre. Há pelo menos três drogas principais que estimulam o centro de prazer no cérebro: dopamina, serotonina e noradrenalina. Todas dão uma sensação de prazer, euforia e falsa confiança, além de produzir um surto de energia.

A cocaína e outros estimulantes funcionam de duas formas. Primeiro, estimulam diretamente as células da ATV para produzirem mais dopamina, provocando um fluxo excessivo da substância no núcleo accumbens. Segundo, evitam que as células da ATV voltem à posição “desligada”, mantendo-as numa produção contínua de dopamina. Além disso, impedem a absorção de serotonina e noradrenalina. O fluxo simultâneo desses três neurotransmissores nos circuitos neurais cria o grande “barato” associado à cocaína.

Em contraste, a heroína e outros opiáceos funcionam neutralizando as células da ATV que podem reduzir a produção de dopamina, gerando assim uma superprodução desta.

Drogas como o LSD operam estimulando a produção de serotonina, que induz uma sensação de bem-estar, determinação e afeição. Mas também estimulam áreas do lobo temporal a alucinações. (Bastam 50 microgramas de LSD para produzir alucinações. A absorção do LSD é tão intensa que aumentar a dose não surte maior efeito.)

Com o passar do tempo, a CIA percebeu que drogas que alteram a mente não eram a solução mágica que estavam procurando. As alucinações e os vícios que acompanham essas drogas as tornam muito instáveis e imprevisíveis, podendo gerar mais problemas do que benefícios em situações políticas delicadas.

Vale notar que somente nos últimos anos a IRM do cérebro de viciados em drogas indicou um novo meio de curar, ou pelo menos tratar, algumas formas de vício. Por acaso, viu-se que vítimas de derrame com lesão na ínsula (localizada no meio do cérebro, entre o córtex pré-frontal e o temporal) tinham uma facilidade maior para deixar de fumar do que a média dos fumantes. Esse resultado foi verificado em usuários de drogas como cocaína, álcool, opiáceos e nicotina. Se esse resultado for confirmado, pode significar que, amortecendo a atividade da ínsula por meio de eletrodos ou estimuladores magnéticos, seja

possível tratar o vício. “É a primeira vez que vemos algo assim, uma lesão numa área cerebral específica que pode eliminar totalmente o problema do vício. É incompreensível”, disse a dra. Nora Volkow, diretora do National Institute on Drug Abuse. No momento, não se sabe como isso funciona, porque a ínsula está ligada a uma grande variedade de funções cerebrais, inclusive percepção, controle motor e consciência de si. Mas se o resultado for comprovado, pode mudar todo o cenário dos estudos sobre vício em drogas.

INVESTIGANDO O CÉREBRO COM OPTOGENÉTICA

Esses experimentos em controle da mente foram realizados principalmente numa época em que o cérebro era um grande mistério, com métodos de tentativa e erro que geralmente fracassavam. Entretanto, devido à explosão de dispositivos de sondagem cerebral, surgiram novas oportunidades que não apenas podem nos ajudar a entender o cérebro, mas também nos ensinar a controlá-lo.

A optogenética, como vimos, é um dos campos científicos que se desenvolve com maior velocidade nos dias de hoje. O objetivo básico é identificar precisamente qual percurso neural corresponde a qual tipo de comportamento. A optogenética começa com um gene chamado opsina, muito peculiar, porque é sensível à luz. (Acredita-se que o surgimento desse gene, centenas de milhões de anos atrás, tenha sido responsável pela criação do primeiro olho. Nessa teoria, um pedacinho de pele sensível à luz evoluiu para se tornar a retina.)

Quando o gene opsina é inserido num neurônio e exposto à luz, o neurônio pode ser acionado. Ligando um interruptor, pode-se reconhecer imediatamente o percurso neural para certos comportamentos, porque as proteínas fabricadas pela opsina conduzem eletricidade e desaparecem.

A parte mais difícil, porém, é inserir esse gene num único neurônio. Para isso, usa-se uma técnica tomada emprestada da engenharia genética. O gene opsina é inserido num vírus inócuo (cujos genes maus foram removidos) e, com instrumentos de precisão, é possível aplicar esse vírus num só neurônio. O vírus infecta o neurônio, inserindo seus genes nele. Quando se irradia luz no tecido neural, esse neurônio é ativado. Desse modo é possível estabelecer o percurso exato de determinadas mensagens.

A optogenética não se limita a permitir a identificação de certos percursos irradiando luz sobre eles, mas também permite aos cientistas controlar o comportamento. Esse método já mostrou ser um sucesso. Há muito tempo se suspeitava de que um simples circuito neural era capaz de espantar as moscas-das-frutas. Usando esse método, foi possível identificar o percurso exato por trás dessa fuga. Bastava acender uma luz nas moscas e elas saíam voando

obedientemente.

Hoje, usando um feixe de luz, os cientistas podem até fazer minhocas pararem de se contorcer e, em 2011, fizeram outra grande descoberta. Em Stanford, eles conseguiram inserir um gene opsina numa região exata da amígdala de um rato. Criados especificamente para serem tímidos, esses ratos ficavam encolhidos na gaiola, mas, quando um feixe de luz foi lançado no cérebro deles, perderam imediatamente a timidez e passaram a examinar a gaiola.

As implicações são enormes. Enquanto as moscas-das-frutas podem ter mecanismos de reflexo simples, envolvendo apenas alguns neurônios, os camundongos possuem um sistema límbico que tem contrapartes no ser humano. Apesar de muitos experimentos que funcionam com camundongos não se aplicarem a seres humanos, é possível que algum dia os cientistas encontrem os caminhos neurais exatos de certas doenças mentais, e que a partir daí seja possível tratá-las sem efeitos colaterais. Como disse o dr. Edward Boyden, do MIT, “Se quisermos desligar um circuito cerebral e a alternativa for a remoção cirúrgica de uma região cerebral, o implante de fibra ótica pode ser o mais indicado”.

Uma aplicação prática é o tratamento do mal de Parkinson. Como vimos, a doença pode ser tratada com estimulação cerebral profunda, mas, como falta precisão ao posicionamento dos eletrodos no cérebro, há sempre o perigo de derrame, hemorragia, infecção etc. A estimulação profunda do cérebro pode também causar efeitos colaterais como tontura e contrações musculares, porque os eletrodos podem, acidentalmente, estimular neurônios indevidos. A optogenética pode aprimorar a estimulação cerebral profunda, ao identificar exatamente os caminhos neurais que não disparam corretamente, no nível de neurônios individuais.

Vítimas de paralisia também podem se beneficiar dessa nova tecnologia. Como vimos no capítulo 4, alguns indivíduos tetraplégicos foram ligados a um computador para controlar um braço mecânico mas, como não têm o sentido do tato, podem derrubar ou quebrar o objeto que desejam pegar. “Ao enviar informação dos sensores, colocados nas próteses dos dedos, diretamente para o cérebro com o uso da optogenética, em princípio, é possível fornecer um sentido do tato bastante preciso”, disse o dr. Krishna Shenoy, de Stanford.

A optogenética pode também ajudar a esclarecer quais são os percursos neurais envolvidos no comportamento humano. De fato, já existem planos de experimentação dessa técnica com seres humanos, principalmente no caso de doenças mentais. É claro que haverá grandes dificuldades. Primeiro, a técnica exige abrir o crânio e, se os neurônios a serem estudados estiverem localizados lá no meio do cérebro, o procedimento será ainda mais invasivo. Além disso, será preciso inserir no cérebro fios muito finos que irradiem luz no neurônio modificado para ativar o comportamento desejado.

Quando esses percursos neurais forem decifrados, será possível estimulá-los, provocando comportamentos estranhos em animais (como, por exemplo, camundongos correndo em círculos). Embora os cientistas estejam apenas começando a traçar os percursos neurais que controlam comportamentos animais simples, no futuro também haverá uma enciclopédia de comportamentos humanos. Contudo, em mãos erradas, a optogenética pode ser usada para exercer controle sobre o comportamento humano.

No geral, os benefícios da optogenética superam em muito suas desvantagens. Ela pode literalmente revelar os percursos neurais a serem usados no tratamento de doenças mentais, e outras. Também pode dar aos cientistas instrumentos para reparar lesões, e talvez curar doenças consideradas incuráveis. No futuro próximo, os resultados serão todos positivos. E num futuro distante, quando todos os percursos do comportamento humano forem entendidos, a optogenética poderá ser usada também para controlar, ou pelo menos modificar, o comportamento humano.

O CONTROLE DA MENTE E O FUTURO

Resumindo, o uso de drogas e de hipnose pela CIA foi um fiasco. São técnicas muito instáveis e imprevisíveis para serem usadas pelo Exército. Podem ser utilizadas para induzir alucinações e dependência, mas não servem para apagar memória, tornar alguém mais manipulável, nem forçar ninguém a praticar atos contra sua vontade. Os governos vão continuar tentando, mas o resultado é duvidoso. Até o momento, as drogas são um instrumento muito agressivo de controle de comportamento.

Há até um conto de advertência: Carl Sagan fala de um cenário assustador e muito possível. Ele imagina um ditador colocando eletrodos nos centros de “dor” e de “prazer” do cérebro de crianças. Os eletrodos são ligados a um computador sem fio, de modo que o ditador pode controlá-las ao apertar um botão.

Outro pesadelo envolve a instalação de sondas no cérebro capazes de anular nossos desejos e assumir o controle muscular, nos obrigando a atos que não queremos fazer. O trabalho de Delgado é incipiente, mas mostra que rajadas de eletricidade aplicadas a áreas motoras do cérebro podem anular nossos pensamentos conscientes, de modo que os músculos fogem ao controle. Ele conseguiu identificar apenas alguns comportamentos de animais que podiam ser controlados por meio de sondas elétricas. No futuro, talvez seja possível encontrar uma variedade maior de comportamentos passíveis de serem controlados eletronicamente ao se apertar um botão.

Se você é a pessoa controlada, a experiência é bem desagradável. Enquanto

you think that you are the owner of your body, your muscles can be activated without your permission and you will act against your will. The electric impulse sent to the brain can be greater than your voluntary impulses sent to the muscles, and you have the impression that someone has taken your body. Your own body becomes a strange object.

In principle, a sad fate of this type can be possible in the future. On the other hand, there are various factors capable of preventing it. First, this technology is still being developed, and it is not known how it will be applied to human behavior, so there is still time to monitor its development, and perhaps create safeguards to avoid its misuse. Second, a dictator can decide to use propaganda and coercion, common methods of population control, which are cheap and effective, instead of installing electrodes in the brains of millions of children, which would be a very expensive and invasive procedure. Third, in democratic societies, there would be a fierce public debate about the perspectives and limitations of this radical technology. Laws would be approved to prevent the abuse of this method, without prejudice to its use to reduce human suffering. In the near future, science will give us new and detailed knowledge about neural pathways in the human brain. It is necessary to distinguish the technologies that can benefit society from those that can control it. The condition for passing these laws is to keep the society well informed.

But the real impact of this technology, I believe, will be liberation, and not enslavement of the mind. These technologies give hope to those imprisoned by mental diseases. Despite the fact that there is still no permanent cure for mental disease, these new technologies have given us a deeper knowledge of the formation and evolution of these disorders. One day, combining genetics, medicine and high technology, we will find a way to control, and even cure, these old diseases.

A recent attempt to explore the new knowledge of the brain is to study historical personalities. Perhaps the discoveries of modern science can help explain the mental states of figures from the past.

One of the most extraordinary figures analyzed so far is Joan of Arc.

Os amantes e os loucos têm a mente tão fervilhante...

O louco, o amante e o poeta

São compostos apenas de imaginação.

– WILLIAM SHAKESPEARE, *SONHO DE UMA NOITE DE VERÃO*

9 ESTADOS ALTERADOS DE CONSCIÊNCIA

Joana d'Arc era uma camponesa analfabeta que dizia ouvir vozes diretamente de Deus, e saiu da obscuridade para levar um exército desmoralizado a vitórias que mudaram o curso de nações, fazendo dela uma das figuras mais fascinantes, envolventes e trágicas da história.

Durante o caos da Guerra dos Cem Anos, enquanto o Norte da França era dizimado pelas tropas inglesas e a monarquia francesa fugia, uma jovem de Orléans chegou dizendo ter recebido instruções divinas para levar o Exército francês à vitória. Não tendo mais nada a perder, Carlos VII autorizou-a a comandar algumas tropas. Para espanto geral, ela conseguiu uma série de triunfos sobre os ingleses. As notícias sobre a garota extraordinária se espalharam rapidamente. Sua fama crescia a cada vitória, até que ela se tornou uma heroína, assumindo o comando militar de toda a França. As tropas francesas, antes à beira do colapso total, conseguiram vitórias decisivas para a coroação do novo rei.

Entretanto, ela foi traída e capturada pelos ingleses. Cientes da ameaça que ela representava, pois era um forte símbolo da própria França e alegava receber orientações de Deus, os ingleses a levaram a um julgamento espetacular. Após um interrogatório com final previamente combinado, Joana foi julgada culpada e queimada na fogueira, aos 19 anos, em 1431.

Nos séculos que se seguiram, foram feitas centenas de tentativas de entender essa extraordinária adolescente. Teria sido profeta, santa, ou louca? Mais recentemente, cientistas recorreram à psiquiatria moderna e à neurociência para explicar a vida de personalidades históricas como Joana d'Arc.

Poucos questionam sua sinceridade quanto à inspiração divina. Mas muitos cientistas escreveram que ela devia sofrer de esquizofrenia, pois ouvia vozes. Outros contestaram esse fato, pois os registros remanescentes de seu julgamento revelam uma pessoa de pensamento e fala racionais. Os ingleses lhe armaram diversas armadilhas teológicas. Perguntaram, por exemplo, se ela estava na graça de Deus. Se ela respondesse sim, seria considerada herege, pois ninguém poderia saber ao certo se estava na graça de Deus. Se respondesse não, seria uma confissão de culpa, e, portanto, uma fraude. De qualquer maneira, ela estava perdida.

Numa resposta que deixou todos boquiabertos, ela disse: “Se eu não estiver, que Deus me coloque lá; se eu estiver, que Deus me guarde.” O notário escreveu no registro: “Os que a estavam interrogando ficaram estupefatos.”

De fato, as transcrições dos interrogatórios são tão impressionantes que George Bernard Shaw colocou traduções literais em sua peça *Santa Joana*.

Mais recentemente, surgiu outra teoria sobre essa mulher excepcional: talvez ela sofresse de epilepsia do lobo temporal. As pessoas nessa condição às vezes têm convulsões, mas outras são atingidas por um curioso efeito colateral, que

ajuda a esclarecer a estrutura das crenças humanas. São pacientes que sofrem de “hiper-religiosidade” e pensam realmente que há um espírito ou presença por trás de tudo. Eventos aleatórios nunca são aleatórios, e têm um profundo significado religioso. Alguns psicólogos especularam que muitos profetas da história sofriam dessa lesão epiléptica no lobo temporal, pois tinham certeza de que falavam com Deus. O neurocientista David Eagleman diz: “Uma parte dos profetas, mártires e líderes da história parecem ter tido epilepsia no lobo temporal. Vejamos Joana d’Arc, a menina de 16 anos que virou o jogo na Guerra dos Cem Anos porque acreditava (e convenceu os soldados franceses) que ouvia vozes de São Miguel Arcanjo, Santa Catarina de Alexandria, Santa Margarida e São Gabriel.”

Esse curioso efeito já havia sido observado em 1892, quando compêndios sobre doença mental mencionaram uma conexão entre “emocionalismo religioso” e epilepsia. A primeira descrição clínica foi feita em 1975, pelo neurologista Norman Geschwind, do Boston Veterans Administration Hospital. Ele observou que epilépticos com impulsos elétricos anormais nos lobos temporais costumavam ter experiências religiosas, e especulou que a tempestade elétrica no cérebro poderia ser a causa dessas obsessões religiosas.

O dr. V. S. Ramachandran estima que entre 30% e 40% dos epilépticos do lobo temporal de quem tratou sofrem de hiper-religiosidade. E observa: “Às vezes é um Deus próprio, às vezes é um sentimento mais difuso de ser unificado com o cosmo. Tudo parece estar repleto de significado. O paciente diz ‘Finalmente entendi tudo mesmo, doutor. Agora realmente entendo Deus. Entendo meu lugar no universo – o esquema cósmico’.”

Ele observa ainda que muitos desses indivíduos são extremamente determinados e convincentes em suas crenças: “Às vezes me pergunto se esses pacientes de epilepsia no lobo temporal têm acesso a outra dimensão da realidade, um desses buracos de minhoca para um universo paralelo. Mas evito dizer isso a meus colegas, para não duvidarem da minha sanidade.” Em experimentos com pacientes de epilepsia no lobo temporal, ele confirmou que esses indivíduos tinham uma forte reação emocional à palavra “Deus”, mas não a palavras neutras. Isso significa que a conexão entre a hiper-religiosidade e essa epilepsia é real, e não fictícia.

O psicólogo Michael Persinger afirma que um certo tipo de estimulação elétrica transcraniana (chamada estimulação magnética transcraniana, ou EMT) pode induzir propositalmente o efeito dessas lesões epilépticas. Se assim for, será possível usar campos magnéticos para alterar crenças religiosas?

Nos estudos do dr. Persinger, uma cobaia coloca um capacete (apelidado de “capacete de Deus”), com um dispositivo que envia magnetismo para determinadas áreas do cérebro. Depois, quando o sujeito é entrevistado, costuma afirmar que esteve com algum grande espírito. David Biello, em artigo na revista

Scientific American, diz: “Durante os três minutos de estimulação, as cobaias traduziram sua percepção do divino em sua própria linguagem cultural – denominando-o Deus, Buda, uma presença benevolente, ou a maravilha do universo.” Já que esse efeito pode ser reproduzido, isso indica que talvez o cérebro seja estruturado de modo a responder a sentimentos religiosos.

Alguns cientistas foram mais longe, especulando a existência de um “gene de Deus”, que predispõe o cérebro à religiosidade. Dado que a maioria das sociedades criou algum tipo de religião, é plausível supor que nossa capacidade de responder a sentimentos religiosos seja programada em nosso genoma. Alguns teóricos da evolução vêm tentando explicar esse fato, afirmando que a religião serviu para aumentar as chances de sobrevivência dos primeiros humanos. A religião deve ter ajudado a unir indivíduos hostis para formar uma tribo coesa, com uma mitologia em comum, o que aumentou as chances de união e sobrevivência da tribo.

Um experimento como o do “capacete de Deus” será capaz de abalar a convicção religiosa de alguém? O aparelho de IRM pode registrar a atividade cerebral de alguém que teve uma experiência religiosa?

Para testar essas ideias, o dr. Mario Beauregard, da Universidade de Montreal, reuniu um grupo de 15 freiras carmelitas que concordaram em pôr a cabeça num aparelho de IRM. Para participar do experimento, todas precisavam “ter tido uma experiência de intensa comunhão com Deus”.

A princípio, o dr. Beauregard nutria a esperança de que as freiras tivessem tido uma experiência mística com Deus, que seria registrada pela IRM. Entretanto, enfiadas no aparelho de IRM, com toneladas de bobinas magnéticas e equipamentos de alta tecnologia em volta, não estavam no lugar ideal para uma revelação divina. O melhor que puderam fazer foi evocar a lembrança de experiências anteriores. “Deus não pode ser convocado quando a gente quer”, disse uma freira.

O resultado final foi confuso e inconclusivo, mas várias regiões do cérebro se acenderam durante o experimento:

- O núcleo caudado, ligado à aprendizagem e possivelmente à paixão. (Talvez as freiras estivessem sentindo o amor incondicional de Deus?)
- A ínsula, que monitora as sensações corporais e emoções sociais. (Talvez as freiras estivessem se sentindo muito próximas umas das outras na busca de Deus?)
- O lobo parietal, que ajuda a processar a consciência espacial. (Talvez elas sentissem estar na presença física de Deus?)

O dr. Beauregard admitiu que tantas áreas cerebrais foram ativadas, com tantas interpretações possíveis, que ele não podia afirmar se a hiper-religiosidade podia ser induzida. Contudo, ficou claro para ele que o sentimento religioso das freiras se refletiu na varredura cerebral.

Mas o experimento abalou a fé das freiras em Deus? Não. Pelo contrário, elas entenderam que Deus colocou aquele “rádio” no cérebro para se comunicarem com Ele.

A conclusão delas foi de que Deus criou os humanos para terem essa capacidade, e o cérebro tem uma antena divina dada por Deus para sentirmos Sua presença. David Biello conclui que “Embora os ateus possam argumentar que encontrar a espiritualidade no cérebro signifique que a religião não passa de um delírio religioso, as freiras ficaram animadas com a varredura justamente pelo motivo oposto: aquilo seria uma confirmação da conexão de Deus com elas”. Beauregard concluiu: “Se um ateu vivencia um certo tipo de experiência, ele a relaciona à magnificência do universo. Se for um cristão, associa a Deus. Quem sabe. Talvez sejam a mesma coisa.”

Da mesma forma, Richard Dawkins, biólogo da Universidade de Oxford e ateu ferrenho, certa vez foi colocado no capacete de Deus para ver se haveria mudança em sua crença.

Não houve.

Conclui-se que, embora a hiper-religiosidade possa ser induzida pela epilepsia no lobo temporal, não há evidências convincentes de que campos magnéticos possam alterar as convicções religiosas de alguém.

DOENÇA MENTAL

Há outro estado alterado de consciência que traz grande sofrimento, tanto para o próprio doente quanto para sua família. É a doença mental. As varreduras cerebrais e a alta tecnologia podem revelar a origem dessa aflição e talvez levar à cura? Se assim for, uma das maiores fontes de sofrimento humano poderá ser eliminada.

Por exemplo: ao longo da história, o tratamento da esquizofrenia era brutal e agressivo. Sintomas típicos dos afetados por esse distúrbio mental degenerativo, que aflige 1% da população, são ouvir vozes imaginárias, sofrer de delírios paranoides e pensamentos desorganizados. Ao longo da história, já foram considerados “possuídos” pelo demônio e expulsos, mortos ou trancafiados. Romances góticos às vezes falam de um parente esquisito, demente, que mora escondido num quarto escuro no porão. Até a Bíblia tem uma passagem em que Jesus encontrou dois seres demoníacos. Os demônios pediram a Jesus que os

levasse a uma criação de porcos. Ele disse “Podem ir”. Quando os demônios entraram no chiqueiro, todos os porcos correram barranco abaixo e morreram afogados no mar.

Mesmo hoje, ainda vemos pessoas com sintomas clássicos de esquizofrenia andando pela rua discutindo consigo mesmas. Em geral, os primeiros sinais surgem no fim da adolescência (nos homens) ou aos vinte e poucos anos (nas mulheres). Alguns esquizofrênicos tiveram uma vida normal e podem inclusive ter realizado grandes feitos até que as vozes se instalassem definitivamente. O caso mais famoso é do ganhador do Prêmio Nobel de Economia de 1994, John Nash, representado por Russell Crowe no filme *Uma mente brilhante*. Na década de 1920, Nash desenvolveu trabalhos pioneiros em economia, teoria dos jogos e matemática pura na Universidade de Princeton. Um orientador dele escreveu uma carta de recomendação com uma única linha: “Este homem é um gênio.” Ele foi capaz de produzir com altíssimo nível intelectual mesmo enquanto era perseguido por delírios. Foi internado aos 31 anos de idade, quando teve um surto maior, e passou muitos anos em instituições, ou viajando pelo mundo, temendo que agentes comunistas o matassem.

Até hoje não há um modo exato, universalmente aceito, de diagnosticar a doença mental. Há esperanças, porém, de que algum dia os cientistas usem varredura cerebral e outras invenções da alta tecnologia para criar instrumentos de diagnóstico confiáveis. Os progressos no tratamento da doença mental têm sido demasiadamente lentos. Após séculos de sofrimento, vítimas da esquizofrenia tiveram um primeiro sinal de alívio nos anos 1950, quando foram descobertas acidentalmente drogas antipsicóticas, como a clorpromazina, capazes de controlar ou, às vezes, até eliminar as vozes que atormentam os doentes mentais.

Acredita-se que essas drogas agem regulando o nível de certos neurotransmissores, como a dopamina. Especificamente, a teoria é de que essas drogas bloqueiam o funcionamento de receptores D2 de certas células nervosas, reduzindo assim o nível de dopamina. A teoria de que as alucinações são em parte causadas pelo excesso de dopamina no sistema límbico e no córtex pré-frontal explica também por que as pessoas que tomam anfetaminas têm alucinações.

A ação da dopamina, por ser essencial às sinapses, tem sido associada também a outros distúrbios. Uma teoria sustenta que a doença de Parkinson é agravada pela falta de dopamina nas sinapses, e a síndrome de Tourette pode ser desencadeada por uma superabundância dela. (As pessoas com síndrome de Tourette têm tiques e movimentos faciais inusitados e uma pequena minoria dispara incontrolavelmente palavras obscenas e comentários profanos, pejorativos.)

Mais recentemente, os cientistas apontaram como outro possível culpado o

nível anormal de glutamato no cérebro. Uma razão para acreditar nisso é que a PCP (a droga chamada pó de anjo) é conhecida por criar alucinações semelhantes às dos esquizofrênicos porque bloqueia o receptor de glutamato chamado NMDA. Uma substância muito promissora no tratamento da esquizofrenia é a clozapina, um remédio relativamente novo que simula a produção de glutamato.

Contudo, as drogas antipsicóticas não são necessariamente a salvação. Em cerca de 20% dos casos, elas inibem os sintomas. Cerca de dois terços dos pacientes encontram algum alívio dos sintomas, mas os demais não são afetados. (Segundo uma teoria, os antipsicóticos imitam uma substância química natural ausente no cérebro dos esquizofrênicos, mas não são uma cópia exata. Assim, é preciso experimentar no paciente vários antipsicóticos, praticamente por tentativa e erro. Além disso, como podem produzir efeitos colaterais desagradáveis, muitos esquizofrênicos interrompem a medicação e têm reincidência.)

Recentemente, varreduras cerebrais de esquizofrênicos realizadas enquanto ouviam vozes ajudaram a explicar esse velho distúrbio. Por exemplo: quando falamos sozinhos em silêncio, certas partes do cérebro se iluminam na IRM, especialmente no lobo temporal (como na área de Wernicke). Quando um esquizofrênico ouve vozes, as mesmas áreas do cérebro se iluminam. O cérebro tem muito trabalho para construir uma narrativa coerente, por isso os esquizofrênicos tentam dar sentido a essas vozes não autorizadas, acreditando que elas se originam de fontes estranhas, como marciais inserindo secretamente pensamentos em seu cérebro. O dr. Michael Sweeney, da Universidade do Estado de Ohio, escreve que: “Os neurônios conectados à sensação de som disparam sozinhos, como um pano encharcado de gasolina se incendia espontaneamente no calor intenso de uma garagem fechada. Sem nada à vista e sem som no ambiente, o cérebro do esquizofrênico cria uma forte ilusão de realidade.”

Geralmente essas vozes parecem vir de outras pessoas, e costumam dar ao paciente ordens banais, mas, às vezes, muito violentas. Enquanto isso, os centros de simulação no córtex pré-frontal parecem um piloto automático e, de certa forma, é como se a consciência do esquizofrênico estivesse fazendo o mesmo tipo de simulações que todos nós fazemos, porém, sem a permissão dele. A pessoa fala literalmente consigo mesma sem saber que é ela quem está falando.

ALUCINAÇÕES

A mente está sempre criando alucinações, mas, na maioria das vezes, são facilmente controladas. Vemos imagens que não existem, ou por exemplo, escutamos falsos sons. Portanto, o córtex cingulado é de importância vital para

diferenciar o real do inventado. Essa parte do cérebro nos ajuda a fazer a distinção entre os estímulos externos e os gerados pela própria mente.

Nos esquizofrênicos, porém, o sistema parece estar lesado e eles não conseguem distinguir as vozes reais das imaginárias. (O córtex cingulado anterior é vital devido à sua localização estratégica, entre o córtex pré-frontal e o sistema límbico. A conexão entre essas duas áreas é uma das mais importantes, porque uma área governa o pensamento racional, e a outra, as emoções.)

Até certo ponto, as alucinações podem ser criadas de propósito. Alucinações ocorrem naturalmente quando o paciente é colocado num quarto totalmente escuro, numa câmara de isolamento, ou num ambiente assustador com ruídos estranhos. Esses são exemplos de que “nossos olhos nos enganam”. Na verdade, o cérebro é que nos engana, criando imagens falsas, tentando dar sentido ao mundo, e apontando ameaças. Esse efeito é chamado de “pareidolia”. Quando vemos as nuvens no céu, vemos imagens de animais, pessoas, personagens de desenhos animados. Isso não depende de nós. Está configurado no cérebro.

Em certo sentido, todas as imagens que vemos, reais e virtuais, são alucinações, porque o cérebro está sempre criando imagens falsas para “fazer sentido”. Como vimos, até as imagens reais são parcialmente falsificadas. Mas, nos doentes mentais, talvez certas regiões do cérebro, como o córtex cingulado anterior, tenham lesões e, por isso, o cérebro confunde realidade e fantasia.

A MENTE OBSESSIVA

Já existem drogas também para curar outro distúrbio mental, o TOC (transtorno obsessivo-compulsivo). Como vimos, a consciência humana implica a mediação de vários mecanismos de feedback. Às vezes, porém, esses mecanismos ficam travados na posição “ligado”.

Um entre cada 40 norte-americanos sofre de TOC. Em casos brandos, o sintoma são compulsões como, por exemplo, voltar várias vezes ao sair de casa para confirmar que trancou a porta. O detetive Adrian Monk, no seriado de TV *Monk*, é um caso brando de TOC. Mas esse distúrbio pode ser tão grave a ponto de alguém, por exemplo, se coçar ou se lavar compulsivamente até sangrar a pele. Alguns pacientes de TOC repetem um comportamento obsessivo durante horas, tornando difícil manter um emprego ou constituir uma família.

Normalmente, alguns tipos de comportamento compulsivo, se moderados, são benignos para nós, pois nos mantêm limpos, saudáveis e em segurança. E foi por essa razão que desenvolvemos esses comportamentos. Mas o paciente com TOC grave não consegue interromper o comportamento, que fica tão exagerado que foge ao controle.

Hoje em dia, varreduras cerebrais revelam como isso acontece, mostrando pelo menos três áreas do cérebro que normalmente nos mantêm saudáveis, mas ficam presas num ciclo de feedback. Primeiro, o córtex orbitofrontal, que vimos no capítulo 1, pode agir como um verificador de fatos, conferindo se trancamos a porta ou lavamos as mãos. Ele nos diz: “Ihh, alguma coisa está errada.” Segundo, o núcleo caudado, situado no gânglio basal, governa as atividades aprendidas que se tornam automáticas, e diz ao corpo: “Faça isso, faça aquilo.” Por fim, o córtex cingulado, que registra as emoções conscientes, inclusive o desconforto, diz: “Ainda me sinto péssimo.”

O professor de psiquiatria Jeffrey Schwartz, da Universidade da Califórnia, em Los Angeles, tentou reunir todas essas informações para explicar como o TOC foge ao controle. Imagine que você sinta uma urgência de lavar as mãos. O córtex orbitofrontal reconhece que alguma coisa não está bem, que suas mãos estão sujas. O núcleo caudado entra em ação e faz você lavar as mãos automaticamente. Então o córtex cingulado registra a satisfação de que suas mãos estão limpas.

Na pessoa com TOC, esse circuito fica alterado. Mesmo depois de constatar que as mãos estão sujas e lavá-las, o indivíduo continua com o sentimento desconfortável de que algo não está bem, que as mãos ainda estão sujas, e fica preso num ciclo de feedback ininterrupto.

Nos anos 1960, percebeu-se que o cloridrato de clomipramina aliviava os sintomas dos pacientes de TOC. Essa e outras drogas desenvolvidas desde então elevam o nível do neurotransmissor serotonina, e em experimentos clínicos mostraram-se capazes de reduzir os sintomas de TOC em até 60%. O dr. Schwartz diz que “o cérebro vai continuar fazendo o que faz, mas você não precisa deixar o cérebro te fazer de bobo”. Certamente, essas drogas não curam, mas trazem alívio aos pacientes de TOC.

TRANSTORNO BIPOLAR

Outra forma comum de doença mental é o transtorno bipolar, em que o paciente tem fases de um otimismo extremo e ilusório, seguidas por períodos de depressão profunda. O transtorno bipolar parece ocorrer em vários membros de uma mesma família e, curiosamente, é muito frequente em artistas. Talvez as grandes obras de arte sejam criadas em surtos de criatividade e otimismo. Há uma lista enorme de famosos com transtorno bipolar, entre celebridades de Hollywood, músicos, artistas plásticos e escritores. A substância lítio pode controlar muitos dos sintomas, mas as causas ainda não são claras.

Uma teoria sustenta que o transtorno bipolar pode ser causado pelo

desequilíbrio entre os hemisférios esquerdo e direito. O dr. Michael Sweeney observou: “As varreduras cerebrais levaram pesquisadores a associar emoções negativas, como a tristeza, ao hemisfério direito, e as positivas, como a alegria, ao esquerdo. Durante pelo menos um século, os neurocientistas observaram uma conexão entre lesões no hemisfério esquerdo e transtornos de humor, inclusive depressão e choro incontrolável. Lesões no hemisfério direito, porém, foram associadas a uma ampla gama de emoções positivas.”

Assim, o hemisfério esquerdo, que é analítico e controla a linguagem, tende a entrar em mania se ficar isolado. O hemisfério direito, pelo contrário, é holístico e tende a frear a mania. V. S. Ramachandran escreveu: “Se ficar sem vigilância, o hemisfério esquerdo provavelmente causará um estado de delírios ou mania... Portanto, parece razoável que haja um ‘advogado do diabo’ no hemisfério direito, que permite a ‘você’ adotar uma visão distanciada, objetiva (alocêntrica), de si mesmo.”

Se a consciência humana implica simulações do futuro, deve computar os resultados de eventos futuros levando em conta certas probabilidades. É preciso, portanto, haver um delicado equilíbrio entre otimismo e pessimismo para estimar as chances de sucesso ou fracasso de determinadas ações.

Mas, em certo sentido, a depressão é o preço a pagar pela capacidade de simular o futuro. Nossa consciência é capaz de imaginar todo tipo de resultados horríveis no futuro e, portanto, está ciente de tudo de ruim que pode acontecer, mesmo que sem uma base realista.

É difícil comprovar muitas dessas teorias, pois as varreduras cerebrais de pessoas clinicamente deprimidas indicam que muitas áreas são afetadas. É difícil apontar a fonte do problema, mas, nos clinicamente deprimidos, a atividade nos lobos parietal e temporal parece estar suprimida, indicando que o paciente foi retirado do mundo externo e está vivendo em seu próprio mundo interno. O córtex ventromedial, em particular, parece ter um papel importante. Essa área, aparentemente, cria uma percepção de que há significado e sentido de totalidade no mundo, e que tudo tem um propósito. A hiperatividade nessa área pode provocar o estado de mania, criando um sentimento de onipotência. A subatividade nessa área é associada à depressão e a um sentimento de que a vida não vale a pena. Assim, é possível que um defeito nessa área seja responsável por oscilações do humor.

UMA TEORIA DA CONSCIÊNCIA E DA DOENÇA MENTAL

A teoria do espaço-tempo da consciência se aplica à doença mental? Como? Essa teoria pode nos dar uma noção mais aprofundada desse distúrbio? Como já

mencionamos, definimos a consciência humana como o processo de criação de um modelo do mundo no espaço e no tempo (principalmente no futuro) por meio da avaliação de muitos ciclos de feedback em vários parâmetros, a fim de atingir um objetivo.

Propusemos que a função principal da consciência humana é simular o futuro, mas essa tarefa não é fácil. O cérebro consegue realizá-la fazendo com que os ciclos de feedback verifiquem e equilibrem uns aos outros. Por exemplo: numa reunião, um diretor competente tenta trazer à tona as discordâncias entre os membros da equipe e definir os pontos de vista divergentes, com o objetivo de reunir os diversos argumentos e então tomar a decisão final. Da mesma forma, várias regiões do cérebro fazem estimativas divergentes quanto ao futuro, que são levadas ao córtex pré-frontal dorsolateral, o diretor-geral do cérebro. Essas avaliações discordantes são analisadas e pesadas até ser tomada uma decisão equilibrada.

Podemos agora aplicar a teoria de espaço-tempo da consciência para encontrarmos uma definição para a maioria das formas de doença mental:

A doença mental é em grande parte causada pela ruptura do sistema delicado de controle e equilíbrio entre ciclos de feedback rivais que simulam o futuro (geralmente porque uma região do cérebro é marcada pela hiperatividade, e a outra pela subatividade).

Como o diretor-geral do cérebro (o córtex pré-frontal dorsolateral) já não consegue mais avaliar os fatos de forma equilibrada devido à ruptura dos ciclos de feedback, ele começa a tirar conclusões estranhas e a agir de maneira esquisita. Essa teoria tem a vantagem de ser testável. É preciso fazer varreduras cerebrais por IRM de um doente mental durante seu comportamento disfuncional, avaliando o desempenho dos ciclos de feedback e comparando-os com varreduras por IRM de pessoas normais. Se a teoria estiver correta, o comportamento disfuncional (por exemplo, ouvir vozes ou ficar obsessivo) poderá ser rastreado até um mau funcionamento do sistema de controle e equilíbrio entre ciclos de feedback. A teoria poderá ser refutada se o comportamento disfuncional for totalmente independente da interação entre essas regiões cerebrais.

Dada essa nova teoria da doença mental, podemos agora aplicá-la a várias formas de distúrbios mentais, resumindo a discussão anterior sob essa nova luz.

Já vimos que o comportamento obsessivo de pessoas que sofrem de TOC pode surgir quando há descompasso no sistema de controle e equilíbrio entre vários ciclos de feedback, um registrando algo como impróprio, outro executando uma ação corretiva, e outro sinalizando que o assunto está sendo resolvido. O fracasso do sistema de controle e equilíbrio nesse circuito pode trancar o cérebro num

círculo vicioso, de modo que a mente não consegue acreditar que o problema foi resolvido.

As vozes ouvidas pelos esquizofrênicos podem surgir quando vários ciclos de feedback já não se contrabalançam mais. Um ciclo de feedback gera vozes falsas no córtex temporal (isto é, o cérebro fala consigo mesmo). Alucinações visuais e auditivas são frequentemente verificadas pelo córtex cingulado anterior, de modo que uma pessoa normal pode distinguir entre vozes reais e fictícias. Se essa região não estiver funcionando adequadamente, o cérebro é inundado por vozes incorpóreas que acredita serem reais. Isso pode causar o comportamento esquizofrênico. Da mesma forma, as oscilações maniaco-depressivas dos bipolares podem remontar a um desequilíbrio entre os hemisférios esquerdo e direito. A interação necessária entre avaliações otimistas e pessimistas se desequilibra e há uma oscilação brusca entre os dois modos de humor.

A paranoia também pode ser entendida de acordo com essa mesma ideia. Resulta do desequilíbrio entre a amígdala (que registra o medo e exagera ameaças) e o córtex pré-frontal, que avalia as ameaças em suas devidas proporções.

É preciso também enfatizar que a evolução nos deu esses ciclos de feedback por uma razão: para nos proteger. Eles nos mantêm limpos, saudáveis e conectados socialmente. O problema ocorre quando a dinâmica entre os circuitos é rompida.

Essa teoria pode ser resumida, de forma geral, assim:

DOENÇA MENTAL: Paranoia

CICLO DE FEEDBACK I: Percebendo uma ameaça

CICLO DE FEEDBACK II: Ignorando ameaças

REGIÃO CEREBRAL AFETADA: Amígdala/lobo pré-frontal

DOENÇA MENTAL: Esquizofrenia

CICLO DE FEEDBACK I: Criando vozes

CICLO DE FEEDBACK II: Ignorando vozes

REGIÃO CEREBRAL AFETADA: Lobo temporal esquerdo/córtex cingulado anterior

DOENÇA MENTAL: Transtorno bipolar

CICLO DE FEEDBACK I: Otimismo

CICLO DE FEEDBACK II: Pessimismo

REGIÃO CEREBRAL AFETADA: Hemisfério esquerdo/direito

DOENÇA MENTAL: TOC

CICLO DE FEEDBACK I: Ansiedade

CICLO DE FEEDBACK II: Satisfação

REGIÃO CEREBRAL AFETADA: Córtex orbitofrontal/núcleo caudado/córtex cingulado

Segundo a teoria do espaço-tempo da consciência, muitas formas de doença mental são caracterizadas pela ruptura do sistema de controle e equilíbrio de ciclos de feedback opostos em regiões do cérebro que simulam o futuro. As varreduras cerebrais vêm identificando gradualmente essas regiões. Um entendimento mais completo da doença mental irá revelar, sem dúvida, o envolvimento de muitas outras regiões cerebrais. Este é apenas um esboço preliminar.

ESTIMULAÇÃO CEREBRAL PROFUNDA

Embora a teoria espaço-tempo da consciência possa nos dar uma noção da origem da doença mental, não nos diz como criar novas terapias e remédios.

Como a ciência irá lidar com a doença mental no futuro? É difícil prever, pois sabemos hoje que a doença mental não se resume a uma categoria, mas é toda uma gama de doenças que afligem a mente de diversas formas. Além disso, a ciência por trás da doença mental ainda está dando os primeiros passos, com áreas imensas totalmente inexploradas e inexplicadas.

Hoje está sendo testado um novo método para tratar a infindável agonia daqueles que sofrem de uma das mais comuns, e mais persistentes formas de doença mental, a depressão, que aflige 20 milhões de pessoas apenas nos Estados Unidos. Dez por cento delas sofrem de uma forma incurável de depressão, que até agora resistiu a todos os avanços da medicina. Uma forma direta de tratá-la é usar sondas em certas regiões profundas do cérebro.

Uma pista importante para desvendar essa doença foi descoberta pela dra. Helen Mayberg e colegas, numa pesquisa realizada na Washington University Medical School. Por meio de varreduras, eles identificaram uma parte do cérebro chamada área de Brodmann 25 (também conhecida como região cingulada subgenual) no córtex, que é consideravelmente hiperativa em indivíduos deprimidos para os quais outras formas de tratamento não tiveram sucesso.

Os cientistas usaram estimulação cerebral profunda (ECP) nessa área, inserindo uma pequena sonda e aplicando um choque elétrico, muito parecido

com um marca-passo. O sucesso da ECP é impressionante no tratamento de vários distúrbios. Na última década, a ECP foi usada em 40 mil pacientes de doenças motoras, como Parkinson e epilepsia, que provocam movimentos corporais incontroláveis. Entre 60% e 100% dos pacientes relatam melhora significativa no controle do tremor das mãos. Mais de 250 hospitais, somente nos Estados Unidos, fazem tratamentos com ECP.

Depois a dra. Mayberg teve a ideia de aplicar a ECP diretamente na área de Brodmann 25 para tratar a depressão. Sua equipe atendeu 12 pacientes clinicamente deprimidos que não tinham apresentado melhora após exaustivo uso de medicamentos, psicoterapia e eletrochoques.

Oito desses 12 pacientes apresentaram progresso imediato. O sucesso foi tão surpreendente que outros laboratórios se esforçaram para expandir os resultados, aplicando a ECP a outros transtornos mentais. No momento, está sendo aplicada em 35 pacientes na Emory University e em 30 pacientes de outras instituições.

Mayberg diz: “A Depressão 1.0 podia ser tratada com psicoterapia – as pessoas discutiam de quem era a culpa. A Depressão 2.0 era tida como um desequilíbrio químico. Esta é a Depressão 3.0. O que agora se acredita é que ao dissecar um distúrbio complexo de comportamento chegando a seus sistemas constituintes, temos uma nova forma de pensar sobre o assunto.”

Embora o sucesso da ECP no tratamento de indivíduos deprimidos seja notável, ainda exige muita pesquisa. Primeiro, não está claro por que a ECP funciona. Supõe-se que destrói ou debilita áreas hiperativas do cérebro (como no Parkinson e a área de Brodmann 25) e, conseqüentemente, é eficaz apenas contra doenças causadas por essa hiperatividade. Segundo, é necessário aprimorar a precisão desse procedimento. Apesar de ter sido usado para tratar vários males do cérebro, como a dor do membro fantasma (quando a pessoa sente dor num membro que foi amputado), a síndrome de Tourette e o transtorno obsessivo-compulsivo, o eletrodo inserido no cérebro não tem precisão, e afeta talvez milhões de neurônios em vez dos poucos que são a fonte do problema.

O tempo irá aperfeiçoar a eficiência dessa terapia. Usando tecnologia microeletromecânica (MEM), é possível criar eletrodos microscópicos capazes de estimular apenas alguns neurônios de cada vez. A nanotecnologia também poderá viabilizar nanossondas neurais da espessura de uma molécula, como nanotubos de carbono. E à medida que a sensibilidade da IRM aumentar, facilitará a colocação desses eletrodos em áreas mais específicas.

DESPERTANDO DE UM COMA

A estimulação cerebral profunda tem se ramificado em diversos campos de

pesquisa, gerando inclusive um efeito colateral benéfico: o aumento do número de células da memória no hipocampo. E é aplicada também para despertar indivíduos em coma.

O coma representa uma das formas talvez mais controversas de consciência, e geralmente resulta em manchetes de alcance nacional. O caso de Terri Schiavo, por exemplo, fascinou o público. Em 1990, devido a um infarto, ela teve falta de oxigênio que provocou uma grave lesão cerebral. Em consequência, Terri entrou em coma. Seu marido, com aprovação dos médicos, quis preservar sua dignidade de morrer em paz. Mas a família dela achou crueldade desligar os aparelhos de uma pessoa que ainda tinha alguma reação a estímulos, e que um dia poderia despertar miraculosamente. Argumentaram que já ocorreram casos sensacionais de pacientes que recuperaram a consciência após anos de vida vegetativa.

Para resolver a questão, usaram varreduras cerebrais. Em 2003, a maioria dos neurologistas que examinaram a tomografia axial computadorizada (TAC) concluiu que a lesão cerebral de Terri era tão extensa que ela jamais poderia despertar, e só lhe restava ficar num estado vegetativo permanente (EVP). Após sua morte, em 2005, a autópsia confirmou que ela não tivera chance de despertar.

No entanto, em outros casos de pacientes em coma a varredura cerebral mostra que a lesão não é tão grave, e há uma pequena chance de recuperação. No verão de 2007, em Cleveland, após uma estimulação cerebral profunda, um homem acordou e cumprimentou a mãe. Ele tinha sofrido uma extensa lesão cerebral oito anos antes e entrado em coma profundo, conhecido como estado de consciência mínima.

O dr. Ali Rezai liderou a equipe de cirurgiões que realizou a operação. Eles inseriram um par de fios no cérebro do paciente até atingir o tálamo, que, como vimos, é o portal onde a informação sensorial é processada inicialmente. Foi enviada uma corrente de baixa voltagem para estimular a área, o que despertou o homem do coma profundo. (Em geral, ao atingir o cérebro, a eletricidade provoca um desligamento da parte atingida, mas, em certas circunstâncias, pode agir como um choque que põe os neurônios em ação.)

Os avanços da tecnologia de ECP devem aumentar o número de casos bem-sucedidos em diversas áreas. Hoje o eletrodo de ECP tem cerca de 1,5 milímetro de diâmetro, mas toca até um milhão de neurônios quando inserido no cérebro, o que pode causar sangramento e dano nos vasos sanguíneos. De 1% a 3% dos pacientes tratados com ECP, de fato, têm sangramentos que podem evoluir para um derrame. A carga elétrica transmitida pelas sondas de ECP ainda é muito bruta, pulsando a uma taxa constante. Algum dia os cirurgiões serão capazes de ajustar a carga elétrica conduzida pelos eletrodos de modo que as sondas poderão se adequar melhor a cada pessoa e doença. A próxima geração de sondas de

ECP está destinada a ser mais precisa e mais segura.

AGENÉTICA DA DOENÇA MENTAL

Outra tentativa de entender e tratar a doença mental é rastrear suas raízes genéticas. Já foram feitas muitas tentativas nessa área, com resultados confusos e decepcionantes. Há evidências consideráveis de que a esquizofrenia e o transtorno bipolar ocorrem em famílias, mas as tentativas de encontrar os genes comuns a todos os indivíduos não foram conclusivas. Ocasionalmente, cientistas estudaram a árvore genealógica de indivíduos que sofrem de doença mental, e encontraram um gene prevalente. Mas as tentativas de generalizar os resultados foram infrutíferas. Os cientistas só puderam concluir que são necessários fatores ambientais e uma combinação de vários genes para desencadear a doença mental. Entretanto, é geralmente aceito que cada distúrbio tem sua própria base genética.

Em 2012, porém, um dos estudos mais abrangentes já realizados mostrou que de fato pode haver um fator genético comum às doenças mentais. Cientistas da Harvard Medical School e do Massachusetts General Hospital analisaram 60 mil pessoas em todo o mundo e viram que havia uma conexão genética entre cinco das principais doenças mentais: esquizofrenia, transtorno bipolar, autismo, depressão grave e déficit de atenção e hiperatividade (DDAH). Juntas, representam uma fração significativa de todos os pacientes de doença mental.

Após uma análise minuciosa do DNA dos pacientes pesquisados, os cientistas concluíram que quatro genes aumentavam o risco de doença mental. Dois deles envolviam a regulação de canais de cálcio nos neurônios (o cálcio é uma substância química essencial no processamento dos sinais neurais). O dr. Jordan Smoller, da Harvard Medical School, diz que “as descobertas sobre os canais de cálcio sugerem que talvez – ênfase no ‘talvez’ – os tratamentos que afetam o funcionamento da canalização de tal elemento podem ter efeitos sobre uma série de distúrbios”. Bloqueadores de canais de cálcio já estão sendo usados para tratar pessoas com transtorno bipolar. No futuro, esses bloqueadores poderão ser usados também no tratamento de outras doenças mentais.

Esse resultado pode ajudar a explicar o curioso fato de que, quando a doença mental tem ocorrências numa família, seus membros podem manifestar diferentes tipos de transtornos. Por exemplo: se um gêmeo tem esquizofrenia, o outro pode ter uma doença totalmente diferente, como o transtorno bipolar.

A questão é que, embora cada doença mental tenha elementos desencadeadores e genes próprios, pode haver um traço comum. Isolar os fatores comuns a essas doenças pode nos ajudar a encontrar os medicamentos

mais eficazes contra elas.

“O que identificamos aqui é provavelmente a ponta do iceberg”, diz o dr. Smoller. “À medida que os estudos se ampliam, esperamos encontrar outros genes, que podem estar sobrepostos.” Se forem encontrados mais genes nessas cinco doenças, poderemos ter acesso a uma abordagem inteiramente nova da doença mental.

Se forem encontrados mais genes em comum, talvez a geneterapia possa restaurar os danos causados pelos genes problemáticos. Ou abrir a possibilidade de criar novas drogas para tratar a doença em nível neural.

CAMINHOS POSSÍVEIS

Atualmente, portanto, não existe cura para os pacientes de doença mental. Historicamente, os médicos são impotentes para tratá-los, mas a medicina moderna nos deu várias possibilidades e terapias para enfrentar esse velho problema. Algumas delas são:

1. Encontrar novos neurotransmissores e drogas para regular os sinais dos neurônios.
2. Localizar genes ligados a várias doenças mentais, e talvez usar geneterapia.
3. Usar estimulação cerebral profunda para amortecer ou aumentar a atividade neural em certas áreas.
4. Usar EEG, IRM, MEG e EMT para entender exatamente o mau funcionamento do cérebro.
5. No capítulo sobre a engenharia reversa do cérebro, vamos explorar mais um caminho promissor, mapeando todo o cérebro e seus percursos neurais. Isso poderá finalmente desvendar o mistério da doença mental.

Mas, para entender a ampla variedade de doenças mentais, alguns cientistas acreditam que elas podem ser reunidas em dois grupos principais, e cada um deles requer uma abordagem diferente:

1. Transtornos mentais envolvendo lesão cerebral.

2. Transtornos mentais desencadeados por conexões incorretas no cérebro.

O primeiro tipo inclui Parkinson, epilepsia, Alzheimer e uma longa série de transtornos causados por derrames e tumores, em que o tecido cerebral está danificado ou defeituoso. No caso do mal de Parkinson e da epilepsia há neurônios hiperativos numa determinada área. No mal de Alzheimer, um acúmulo de placas amiloides destrói tecidos cerebrais, inclusive no hipocampo. Em derrames e tumores, certas partes do cérebro são silenciadas, causando diversos problemas comportamentais. Cada um desses distúrbios deve ser tratado de maneira diferente, pois cada um deriva de uma lesão distinta. Para o mal de Parkinson e a epilepsia, é indicado o uso de sondas para silenciar as áreas hiperativas, ao passo que os danos provocados por mal de Alzheimer, derrames e tumores são geralmente incuráveis.

No futuro, teremos avanços em outros métodos, além da estimulação profunda e dos campos magnéticos, para tratar essas partes danificadas do cérebro. Um dia, células-tronco poderão substituir tecidos cerebrais danificados. Ou talvez, computadores possam criar partes de reposição artificiais para compensar as áreas danificadas. Nesse caso, o tecido danificado é removido ou substituído orgânica ou eletronicamente.

A segunda categoria engloba distúrbios causados por problemas nas conexões do cérebro. Doenças como esquizofrenia, TOC, depressão e transtorno bipolar pertencem a essa categoria. Cada região do cérebro pode estar relativamente saudável e intacta, mas uma ou mais delas podem estar mal conectadas, fazendo com que as mensagens sejam processadas incorretamente. É difícil tratar essa categoria, dado que as conexões cerebrais não são bem compreendidas. Até o momento, o principal meio de lidar com esses transtornos é o uso de medicamentos que influenciam os neurotransmissores, mas essa abordagem ainda se baseia muito em tentativa e erro.

Existe ainda outro estado alterado de consciência que tem nos ajudado a compreender a mente em ação, abrindo novas perspectivas sobre o funcionamento do cérebro e sobre o que pode acontecer quando há um distúrbio. É o campo da IA, a Inteligência Artificial. Apesar de ainda estar engatinhando, já proporcionou uma visão mais profunda do processo de pensamento e ampliou nosso conhecimento da consciência humana. Isso leva às seguintes questões: É possível criar uma consciência de silício? Se sim, como poderá diferir da consciência humana? E tentará nos controlar?

Não. Não estou interessado em desenvolver um cérebro potente. Só estou querendo um cérebro medíocre, algo como o do presidente da American Telephone and Telegraph Company.

– ALAN TURING

10 A MENTE ARTIFICIAL E A CONSCIÊNCIA DE SILÍCIO

Em fevereiro de 2011, um acontecimento entrou para a história.

Um computador da IBM chamado Watson fez o que muitos críticos consideravam impossível: venceu dois competidores num programa de perguntas e prêmios na TV, chamado *Jeopardy!*. Milhões de telespectadores ficaram grudados na tela enquanto Watson eliminava metodicamente os outros candidatos em cadeia nacional, respondendo a perguntas que deixaram seus rivais embasbacados, e fez jus ao prêmio de um milhão de dólares.

A IBM superou todas as limitações, fabricando uma máquina com potência monumental. Watson é capaz de processar dados com a espantosa velocidade de 500 gigabytes por segundo (o equivalente a um milhão de livros por segundo), com 16 trilhões de bytes de memória RAM. Conseguiu acessar 200 milhões de páginas de material em sua memória, inclusive todo o conhecimento armazenado na Wikipédia, e foi capaz de analisar essa montanha de informações num programa de TV ao vivo.

Watson é a última geração de “sistemas especializados”, programas de computador que usam a lógica formal para acessar grandes quantidades de informação especializada (uma máquina que atende ao telefone e lhe dá um menu de opções é um sistema especializado muito primitivo). Os sistemas especializados vão continuar a evoluir, tornando nossa vida mais conveniente e eficiente.

Por exemplo: atualmente, engenheiros estão trabalhando na criação de um “robo-doc” [robô-médico], que vai aparecer em nosso relógio de pulso ou numa tela para nos dar orientação médica com 99% de exatidão, quase de graça. Falamos os sintomas, e ele acessa os bancos de dados dos melhores centros médicos do mundo inteiro, buscando as últimas informações científicas. Isso vai reduzir as idas desnecessárias ao consultório, eliminar os falsos alarmes, e facilitar as consultas médicas periódicas.

Um dia poderemos ter robôs-advogados, que saibam responder a todas as questões legais corriqueiras, e robôs-secretários para planejar férias, fazer reservas de viagens e restaurantes. É claro que, para serviços mais especializados, que exigem orientação profissional, ainda precisaremos de um médico de verdade, um advogado de verdade etc., mas para aconselhamento de questões banais do cotidiano, esses programas serão muito convenientes.

Além disso, os cientistas criaram “chat-bots” [conversas com robôs] que imitam conversas comuns. Qualquer pessoa conhece, em média, dezenas de milhares de palavras. Para ler um jornal, é preciso conhecer duas mil palavras ou mais; no entanto, uma conversa casual geralmente exige poucas centenas. Há robôs que podem ser programados para conversar dentro desse vocabulário reduzido (desde que a conversa se limite a certos temas bem definidos).

FUROR NA MÍDIA – OS ROBÔS ESTÃO CHEGANDO

Logo após Watson ter vencido a competição, alguns formadores de opinião ficaram agitados, já lamentando o dia em que as máquinas irão assumir o controle. Ken Jennings, um dos competidores derrotados por Watson, disse à imprensa: “De minha parte, dou boas-vindas a nossos senhores computadores.” Os jornalistas perguntaram: Se Watson era capaz de vencer participantes veteranos desses programas de TV numa competição entre homem e máquina, que chance teremos nós, reles mortais, de enfrentar as máquinas? Em tom de brincadeira, Jennings respondeu: “Brad (o outro competidor) e eu fomos os primeiros operários da indústria do conhecimento demitidos pela nova geração de máquinas ‘pensantes’.”

Os jornalistas, porém, esqueceram-se de dizer que não era possível entrevistar Watson e parabenizá-lo pela vitória. Não podemos lhe dar tapinhas nas costas, nem brindar com ele. Watson não saberia o que isso significa e não fazia a menor ideia de que tinha vencido. Furor da mídia à parte, a verdade é que Watson é uma máquina de calcular altamente sofisticada, capaz de somar (ou pesquisar arquivos de dados) com uma velocidade bilhões de vezes maior que o cérebro humano, mas carece totalmente de consciência de si ou de senso comum.

Por um lado, o progresso no campo da inteligência artificial tem sido espantoso, principalmente na área da potência computacional pura. No ano de 1900, se alguém visse o poder de cálculo dos computadores de hoje, diria se tratar de um milagre. Mas, em outro sentido, o progresso tem sido extremamente lento na construção de máquinas que pensem por si mesmas (isto é, verdadeiros autômatos, que não precisem de um manipulador de marionetes, um controlador com um *joystick*, ou alguém no painel de controle remoto). Os robôs não têm a menor noção de que são robôs.

Dado que a potência dos computadores tem sido duplicada a cada dois anos nos últimos cinquenta anos, pela lei de Moore, alguns dizem que é só uma questão de tempo para que as máquinas adquiram uma autoconsciência capaz de rivalizar com a inteligência humana. Ninguém sabe quando isso irá acontecer, mas a humanidade deve estar preparada para o momento em que a consciência da máquina saia do laboratório e entre no mundo real. O modo de lidar com a consciência dos robôs poderá decidir o futuro da raça humana.

CICLOS DE “ALTOS E BAIXOS” NA INTELIGÊNCIA ARTIFICIAL

É difícil prever o destino da IA, que já passou por três ciclos de altos e baixos. Nos anos 1950, parecia que as empregadas domésticas e os mordomos mecânicos logo se tornariam realidade. Foram construídas máquinas que jogavam xadrez e solucionavam problemas de álgebra. Braços robóticos reconheciam e pegavam blocos. A Universidade de Stanford criou um robô chamado Shakey – basicamente um computador sobre rodas, com uma câmera – que se movimentava sozinho numa sala, evitando obstáculos.

Logo foram publicados artigos precipitados em revistas científicas, anunciando a chegada do robô doméstico. Algumas previsões eram muito conservadoras. Em 1949, a revista *Popular Mechanics* afirmava que “no futuro, os computadores não pesarão mais do que 1,5 tonelada”. Outras eram bastante otimistas, anunciando que o dia dos robôs estava próximo. Um dia, Shakey será um mordomo ou empregada doméstica mecânica, passando aspirador de pó nos tapetes e abrindo as portas para nós. Filmes como *2001: Uma odisseia no espaço* nos convenceram de que os robôs em breve estariam pilotando foguetes espaciais da Terra para Júpiter, e batendo papo com os astronautas. Em 1965, o dr. Herbert Simon, um dos fundadores da IA, disse tranquilamente: “Dentro de 20 anos, as máquinas serão capazes de fazer qualquer trabalho que o homem faz”. Dois anos depois, outro fundador da IA, o dr. Marvin Minsky, disse que “dentro de uma geração (...) a dificuldade de se criar ‘inteligência artificial’ estará praticamente superada”.

Mas todo esse otimismo chegou ao fim nos anos 1970. Robôs jogadores de xadrez só sabiam jogar xadrez. Braços mecânicos só sabiam pegar blocos, e nada mais. Eram como animais treinados para fazer um único número. O mais avançado dos robôs levava horas para atravessar uma sala. Colocado num ambiente estranho, Shakey se perdia com a maior facilidade. Além disso, os cientistas estavam longe de entender a consciência. Em 1974, a IA sofreu um grande golpe quando os Estados Unidos e a Inglaterra reduziram drasticamente os subsídios aos trabalhos na área.

Mas, como a potência dos computadores continuava crescendo nos anos 1980, houve uma nova corrida do ouro pela IA, impulsionada principalmente pela esperança do Pentágono de colocar robôs soldados no campo de batalha. A verba destinada às pesquisas em IA chegou a um bilhão de dólares em 1985, com centenas de milhões de dólares gastos em projetos como o Smart Truck, que seria um caminhão inteligente, autônomo, para penetrar nas linhas inimigas, fazer o reconhecimento do terreno, realizar missões (como resgate de prisioneiros) e retornar ao território aliado. Infelizmente, o único feito realizado pelo caminhão foi se perder. Os notórios fracassos desses projetos caríssimos criaram mais um período de estagnação para a Inteligência Artificial nos anos 1990.

Paul Abrahams, falando sobre a época em que era aluno no MIT, disse: “É

como se um grupo de pessoas se propusesse a construir uma torre para chegar à Lua. A cada ano, eles mostram com orgulho o quanto a torre está mais alta do que no ano anterior. O único problema é que a Lua não está chegando mais perto.”

Mas hoje, com o desenvolvimento incessante da potência dos computadores, teve início mais um renascimento da IA, e estão sendo feitos progressos lentos, porém, substanciais. Em 1997, o computador Deep Blue da IBM venceu o campeão de xadrez Garry Kasparov. Em 2005, o carro robô de Stanford venceu a Darpa Grand Challenge, competição de carros sem motorista. Outras grandes metas têm sido alcançadas.

Resta a questão: A sorte estará na terceira tentativa?

Os cientistas percebem hoje que subestimaram demais o problema porque a maior parte do pensamento humano é inconsciente. De fato, a parte consciente dos nossos pensamentos representa apenas uma parcela mínima de nossas computações.

O dr. Steve Pinker diz: “Eu pagaria um dinheirão por um robô que lavasse os pratos e realizasse tarefas simples, mas não posso, porque todos os probleminhas que é preciso resolver para construir um robô que faça coisas desse tipo, como reconhecer objetos, pensar sobre o mundo e controlar mãos e pés, são problemas de engenharia não resolvidos.”

Embora os filmes de Hollywood nos digam que robôs assustadores como o Exterminador estão logo ali na esquina, a tarefa de criar uma mente artificial tem se mostrado muito mais difícil do que se pensava. Uma vez perguntei ao dr. Minsky quando as máquinas irão se equiparar, ou talvez superar, à inteligência humana. Ele disse estar confiante em que isso irá acontecer, mas não faz mais previsões de datas. Em vista da montanha russa que é a história da IA, talvez o mais sensato seja mapear o futuro da IA sem fazer um cronograma.

RECONHECIMENTO DE PADRÕES E SENSO COMUM

Há pelo menos dois problemas básicos para a IA: o reconhecimento de padrões e o senso comum.

Os melhores robôs mal conseguem reconhecer objetos simples como uma bola ou um copo. O olho do robô pode ver mais detalhes do que um olho natural, mas seu cérebro não reconhece o que vê. Se um robô for colocado numa rua movimentada que ele não conhece, ele se desorienta rapidamente e fica perdido. Devido a esse problema, o reconhecimento de padrões (por exemplo, a identificação de objetos) progrediu muito mais lentamente do que se supunha.

Quando um robô entra numa sala, tem que fazer trilhões de cálculos,

decompondo os objetos que ele vê em pixels, linhas, círculos, quadrados e triângulos, para fazer correspondências com as milhares de imagens armazenadas em sua memória. Por exemplo: o robô vê uma cadeira como uma mistura de linhas e pontos, mas não é capaz de identificar facilmente a essência do que é uma cadeira. Mesmo que consiga ligar um objeto a uma imagem em sua base de dados, uma pequena rotação (como uma cadeira caída no chão) ou uma mudança de perspectiva (a cadeira num ângulo diferente) vai confundir o robô. Automaticamente, nosso cérebro leva em conta diferentes perspectivas e variações. Nosso cérebro realiza inconscientemente trilhões de cálculos, mas o processo não parece exigir esforço algum.

Os robôs têm também dificuldades com o senso comum. Eles não entendem fatos simples do mundo físico e biológico. Não existe uma equação que confirme algo tão óbvio (para nós, humanos) como “o calor excessivo é desconfortável”, ou “as mães são mais velhas que as filhas”. Já houve até algum progresso na tradução desse tipo de informação para a lógica matemática, mas catalogar o senso comum de uma criança de 4 anos exigiria centenas de milhões de linhas de códigos de programação. Como disse Voltaire: “O senso comum não é muito comum.”

Por exemplo: um dos robôs mais avançados, Asimo, foi construído pela Honda Corporation no Japão (onde são feitos 30% de todos os robôs industriais). É um robô maravilhoso, do tamanho de um menino, que anda, corre, sobe escadas, fala várias línguas e dança (muito melhor que eu). Interagi várias vezes com Asimo na televisão, e fiquei impressionado com suas habilidades.

Entretanto, numa reunião particular com os criadores do robô, fiz a pergunta-chave: A inteligência de Asimo é comparável à de qual animal? Eles admitiram que era a inteligência de um inseto. Todo o show de dança e fala era mais destinado à mídia. O problema é que Asimo é, de um modo geral, um grande gravador. Possui apenas uma lista modesta de funções autônomas, de modo que cada movimento ou fala tem que ser detalhadamente programado com antecedência. Por exemplo: um vídeo curto da minha interação com Asimo levou cerca de três horas para ser gravado, porque os gestos das mãos e outros movimentos tinham que ser programados por uma equipe de operadores.

Considerando nossa definição de consciência humana, os robôs atuais parecem estagnados num nível muito primitivo, tentando entender o mundo físico e social por meio da aprendizagem de fatos básicos. Sendo assim, os robôs nem chegaram ao estágio de fazer simulações realistas do futuro. Dizer a um robô que planeje um assalto a banco, por exemplo, implica supor que ele sabe tudo o que é fundamental num banco, como o sistema de segurança, o lugar em que o dinheiro fica guardado, e como a polícia e as pessoas presentes reagirão ao assalto. Alguns desses dados podem ser programados, mas há centenas de nuances que a mente humana entende naturalmente, mas os robôs, não.

Os robôs são excelentes na simulação do futuro apenas em determinadas áreas, como jogar xadrez, fazer cálculos para previsão meteorológica, traçar colisão de galáxias etc. Como as regras do xadrez e a lei da gravidade são conhecidas há séculos, simular o futuro de um jogo ou de um sistema solar é só uma questão de potência computacional.

Tentativas de superar esses limites à força também fracassaram. Um programa ambicioso chamado CYC foi desenvolvido para resolver o problema do senso comum. O CYC continha milhões de linhas de códigos de programação, com todas as informações e conhecimentos necessários para compreender o ambiente físico e social. Mesmo sendo capaz de processar centenas de milhares de fatos e milhões de comandos, o CYC não reproduz o nível de pensamento de um humano de 4 anos de idade. Infelizmente, após algumas notícias otimistas na imprensa, o projeto ficou estagnado. Muitos programadores desistiram, manchetes de jornal apareceram e sumiram, mas ainda assim o projeto não morreu.

O CÉREBRO É UM COMPUTADOR?

Onde foi que erramos? Nos últimos 50 anos, os cientistas que trabalham com Inteligência Artificial vêm tentando definir padrões do cérebro seguindo uma analogia com computadores. Mas talvez isso seja muito simplista. Como disse Joseph Campbell: “Os computadores são como os deuses do Velho Testamento: muitas leis e nenhuma misericórdia.” Retirando um único transistor de um chip Pentium, o computador para imediatamente. Mas o cérebro humano pode continuar funcionando mesmo sem a metade dele.

Isso acontece porque o cérebro não é um computador, e sim um tipo de rede neural altamente sofisticada. Diferentemente de um computador, que tem uma arquitetura fixa (entrada, saída e processador), as redes neurais são conjuntos de neurônios que se reconectam e se reforçam a cada nova tarefa aprendida. O cérebro não tem programação, nem sistema operacional, nem Windows, nem processador central. As redes neurais são basicamente paralelas, com 100 bilhões de neurônios disparando ao mesmo tempo a fim de atingir um único objetivo: aprender.

Com isso em mente, os pesquisadores de IA estão reexaminando a abordagem “de cima para baixo”, que vinham seguindo nos últimos 50 anos (por exemplo, colocando todas as regras de senso comum num CD). Agora eles estão adotando a abordagem “de baixo para cima”. Essa abordagem tenta seguir a Mãe Natureza, que criou seres inteligentes (nós) através da evolução, começando com animais simples, como vermes e peixes, e depois criando outros mais complexos.

As redes neurais aprendem pelo modo mais difícil, tropeçando e cometendo erros.

O dr. Rodney Brooks, ex-diretor do famoso MIT Artificial Intelligence Laboratory e cofundador do iRobot, que fabrica aqueles aspiradores de pó automáticos presentes em muitas casas, introduziu uma abordagem inteiramente nova da IA. Em vez de projetar robôs grandes e desajeitados, por que não fazer robôs pequenos e compactos como insetos, que precisam aprender a andar, tal como acontece na natureza? Quando o entrevistei, o dr. Rodney me disse que ficava maravilhado ao ver que um mosquito, de cérebro quase microscópico com pouquíssimos neurônios, faz manobras no espaço muito melhor que qualquer avião-robô. Ele construiu uma série de robôs extremamente simples, carinhosamente chamados de “insetoides” ou “bugbots”, que corriam pelo MIT e em volta dos robôs mais tradicionais. Seu objetivo era criar robôs que seguissem o método de tentativa e erro da Mãe Natureza. Em outras palavras, os robozinhos aprendiam à medida que esbarravam nas coisas.

A princípio, pode parecer que isso exige muita programação. A ironia, porém, é que a rede neural não exige programação nenhuma. A única coisa que a rede neural faz é se reprogramar, dando menos peso a certos percursos cada vez que toma uma decisão certa. Portanto, programar não é nada; mudar a rede é tudo.

Autores de ficção científica já imaginaram robôs em Marte como humanoides, andando e se movimentando exatamente como nós, dotados de inteligência humana por meio de uma programação complexa. Aconteceu o oposto. Hoje, os netos dessa abordagem – como o jipe-robô Curiosity – estão circulando pela superfície de Marte. Não são programados para andar como humanos, mas, tendo a inteligência de um inseto, se saem muito bem naquele terreno. Esses veículos feitos para Marte têm uma programação relativamente pequena, e aprendem à medida que esbarram em obstáculos.

OS ROBÔS SÃO CONSCIENTES?

Talvez a maneira mais clara de constatar que robôs inteiramente autômatos ainda não existem seja classificar seu grau de consciência. Como vimos no capítulo 2, podemos classificar a consciência em quatro níveis. O Nível 0 refere-se a termostatos e plantas, isto é, envolve alguns ciclos de feedback e poucos parâmetros simples, como temperatura e luz solar. O Nível I descreve insetos e répteis, que têm mobilidade e sistema nervoso central; isso significa criar um modelo de seu mundo em relação a um novo parâmetro, o espaço. O Nível II de consciência cria um modelo de mundo em relação aos outros da mesma espécie, exigindo emoções. O Nível III de consciência descreve os humanos, que

incorporam o tempo e a consciência de si para simular a evolução das coisas no futuro e determinar nosso próprio lugar nesses modelos.

Podemos usar essa teoria para classificar os robôs de hoje. A primeira geração de robôs estava no Nível 0, pois eram estáticos, sem rodas nem esteiras. Os robôs de hoje estão no Nível I, pois são móveis, mas se situam numa escala muito baixa, porque têm uma dificuldade tremenda de se movimentar no mundo real. Sua consciência pode ser comparada à de um verme, ou de um inseto lento. Para preencher totalmente o Nível I, os cientistas terão que criar robôs capazes de duplicar realisticamente a consciência de insetos e répteis. Até os insetos têm capacidades que faltam aos robôs atuais, como encontrar esconderijos rapidamente, localizar parceiros sexuais na floresta, reconhecer predadores e fugir deles, encontrar abrigo e comida.

Como já mencionamos, podemos classificar numericamente a consciência pelo número de ciclos de feedback em cada nível. Os robôs que enxergam, por exemplo, podem ter vários ciclos de feedback porque têm sensores visuais que detectam sombras, bordas, curvas, formas geométricas etc., de forma tridimensional. Da mesma maneira, robôs que escutam têm sensores que detectam frequência, intensidade, ênfase, pausas etc. O número total de seus ciclos de feedback é por volta de 10 (enquanto um inseto, que se alimenta no mato, encontra parceiros, localiza abrigo etc., pode ter 50 ou mais ciclos de feedback). Um robô típico, portanto, deve ter um Nível I:10 de consciência.

Os robôs terão que ser capazes de criar um modelo do mundo em relação aos outros para entrarem no Nível II de consciência. Como já mencionamos, numa primeira aproximação, o Nível II é calculado pela multiplicação da quantidade de integrantes do grupo pelo número de emoções e gestos usados para a comunicação entre eles. Assim, os robôs terão Nível II:0 de consciência. Espera-se que os robôs com emoções, em construção nos laboratórios de hoje, logo subam de nível.

Os robôs atuais veem os humanos como um conjunto de pixels se movendo em seus sensores de TV, mas alguns pesquisadores de IA estão começando a criar robôs capazes de reconhecer emoções nas expressões faciais e no tom de voz dos humanos. É um primeiro passo na direção de robôs que percebam que os humanos são mais do que que pixels aleatórios, e que têm emoções.

Nas próximas décadas, os robôs irão ascender ao Nível II de consciência, tornando-se gradualmente tão inteligentes quanto um camundongo, rato, coelho, até chegar ao gato. Talvez ainda neste século eles sejam tão inteligentes quanto um macaco, e então começarão a criar metas próprias.

Quando os robôs tiverem um conhecimento funcional do senso comum e da Teoria da Mente, serão capazes de fazer simulações complexas do futuro, colocando-se como atores principais, e então entrarão no Nível III de consciência. Sairão do mundo do presente e entrarão no mundo do futuro. Isso

está décadas à frente da capacidade de qualquer robô de hoje. Fazer simulações significa ter uma firme compreensão das leis da natureza, da causalidade e do senso comum, para ser capaz de prever eventos futuros. Significa também entender as intenções e motivações humanas a fim de poder prever o comportamento futuro das pessoas.

O valor numérico do Nível III de consciência, como dissemos, é calculado pelo número total de conexões causais possíveis para simular o futuro em várias situações da vida real dividido pelo valor médio de um grupo de controle. Os computadores de hoje são capazes de fazer simulações limitadas com poucos parâmetros (por exemplo, colisão de duas galáxias, fluxo de ar em torno de um avião, tremor de prédios num terremoto), mas são totalmente despreparados para simular o futuro em situações complexas da vida real, portanto, seu nível de consciência é algo como Nível III:5.

Como podemos ver, talvez leve muitas décadas de trabalho intenso para termos um robô capaz de operar com naturalidade na sociedade humana.

QUEBRA-MOLAS NO CAMINHO

Afinal, quando os robôs terão inteligência igual ou maior que a humana? Não se sabe, mas há muitas previsões. A maior parte se baseia na lei de Moore e aponta para daqui a décadas. Mas a lei de Moore não é uma lei e, na verdade, viola uma lei fundamental da física: a teoria quântica.

Assim, a lei de Moore não pode durar para sempre. Já podemos vê-la enfraquecendo, e pode ficar para trás no fim desta ou da próxima década. Nesse caso, as consequências serão graves, especialmente para o Vale do Silício.

O problema é simples. No momento, é possível colocar centenas de milhões de transistores de silício num chip do tamanho de uma unha, mas há um limite de espaço nesses chips. Hoje, a menor camada de silício num chip Pentium é de aproximadamente vinte átomos de espessura; em 2020, poderá ter cinco átomos a mais. Porém, o princípio de incerteza de Heisenberg entra em cena, e não se pode determinar precisamente onde o elétron está, e ele poderia “vazar” do fio. (Veja o Apêndice, onde discutimos em mais detalhes a teoria quântica e o princípio da incerteza.) O chip entraria em curto-circuito. Além disso, iria gerar calor suficiente para fritar um ovo. Assim, o vazamento e a temperatura irão acabar derrubando a lei de Moore e será necessário fazer logo uma substituição.

Como a acumulação de transistores em chips planos está chegando ao máximo da potência computacional, a Intel está fazendo uma aposta multibilionária de que os chips chegarão à terceira dimensão. O tempo dirá se o jogo vai valer a pena (um problema sério com chips de 3D é que a temperatura aumenta

rapidamente com a altura do chip).

A Microsoft busca outras opções, como expandir em 2D com processamento paralelo. Uma possibilidade é colocar os chips horizontalmente, em fila. Depois, diseca-se o problema de software em partes, cada parte é atribuída a um pequeno chip, e no final as partes são reagrupadas. Entretanto, pode ser um processo difícil, e os softwares se desenvolvem a um ritmo muito mais lento do que a taxa exponencial prevista pela lei de Moore.

Essas medidas paliativas podem acrescentar anos à lei de Moore. Mas em algum momento tudo isso também passará: a teoria quântica assume o comando inevitavelmente. Isso significa que os físicos estão pesquisando uma ampla variedade de alternativas para quando a Idade do Silício estiver chegando ao fim, tais como computadores quânticos, computadores moleculares, nanocomputadores, computadores de DNA, computadores óticos etc. Nenhuma dessas tecnologias, porém, está pronta para o lançamento.

O VALE ESTRANHO

Suponhamos que um dia iremos conviver com robôs incrivelmente sofisticados, talvez usando chips com transistores moleculares em vez de silício. O quanto queremos que esses robôs se pareçam conosco? O Japão é o líder mundial na criação de robôs parecidos com bichinhos e crianças, mas seus projetistas têm o cuidado de não fazê-los parecidos demais com humanos, o que pode causar incômodo. Esse fenômeno foi estudado no Japão pelo dr. Masahiro Mori, nos anos 1970, e é chamado de “vale estranho”. Ele diz que robôs parecidos demais com os humanos ficam assustadores. (Na verdade, esse efeito foi mencionado pela primeira vez por Darwin, em 1839, em *A viagem do Beagle*, e mais tarde por Freud, em 1919, num ensaio intitulado “O estranho”). Desde então, tem sido estudado minuciosamente, não só por pesquisadores de IA, mas também por profissionais de animação, publicitários, e por todos aqueles que promovem produtos envolvendo figuras de aparência humana. Por exemplo: numa resenha do filme *O expresso polar*, um escritor da CNN observou que “Aqueles personagens humanos no filme parecem... bem, assustadores. Assim, *O expresso polar* é na melhor das hipóteses desconcertante, e, na pior, um tanto horripilante.”

Segundo o dr. Mori, quanto mais um robô é parecido com um humano, mais empatia sentimos com relação a ele, mas só até certo ponto. Há uma pequena queda de empatia quando um robô fica com aparência muito humana – daí o “vale estranho”. Um robô muito parecido conosco, mas ainda com alguns traços “estranhos”, cria um sentimento de repulsa e medo. Se um robô parece 100% humano, indistinguível de você e de mim, voltamos a registrar emoções positivas.

Isso tem implicações práticas. Por exemplo: os robôs devem sorrir? Em princípio, parece óbvio que devem sorrir ao cumprimentar as pessoas, fazendo-as se sentirem confortáveis. O sorriso é uma expressão facial universal, que sinaliza acolhimento e boas-vindas. Mas se o sorriso do robô é realista demais, provoca arrepios. (Por exemplo: são comuns as máscaras de Halloween de seres fantasmagóricos sorridentes, com expressão diabólica.) Os robôs só devem sorrir se tiverem traços infantis (isto é, olhos grandes e rostinho redondo), ou perfeitamente humanos, mas nenhuma forma intermediária. Quando o sorriso é forçado, ativamos músculos faciais com o córtex pré-frontal. Quando sorrimos porque estamos de bom humor, os nervos são controlados pelo sistema límbico, que aciona um conjunto ligeiramente diferente de músculos. O cérebro sabe distinguir a sutil diferença entre os dois, o que foi benéfico para nossa evolução.

Esse efeito pode ser estudado também com varreduras cerebrais. Digamos que um indivíduo seja colocado num aparelho de IRM, vendo a imagem de um robô que parece perfeitamente humano, exceto por movimentos bruscos e mecânicos. Cada vez que o cérebro vê alguma coisa, tenta prever o movimento futuro daquele objeto. Quando vemos um robô que parece humano, o cérebro prevê que ele se movimentará como um humano, e se o robô se move como uma máquina, há uma falta de correspondência que nos traz desconforto. O lobo parietal, principalmente, se ilumina (especificamente a parte do lobo em que o córtex motor se conecta ao córtex visual). Acredita-se que existam neurônios espelho nessa área do lobo parietal. Faz sentido, porque o córtex visual capta a imagem do robô semelhante a um humano e seus movimentos são previstos pelo córtex motor e por neurônios espelho. Por fim, é provável que o córtex orbitofrontal, localizado logo atrás dos olhos, coloque tudo junto e diga: “Hummm, tem alguma coisa esquisita aí.”

Os diretores de Hollywood estão cientes desse efeito. Quando gastam milhões para fazer um filme de terror, eles sabem que a cena mais apavorante não é a da bolha gigante atacando ou um enorme vulto de Frankenstein saindo de trás da moita. A cena mais apavorante é quando há uma perversão do comum. Vejamos o filme *O exorcista*. Que cena fez espectadores saírem correndo do cinema para vomitar, ou desmaiarem ali mesmo, na plateia? Foi a cena em que o demônio aparece? Não. Espectadores do mundo inteiro gritaram de terror ou choraram de soluçar quando a cabeça de Linda Blair fez um giro completo.

Essa reação pode ser demonstrada com macacos jovens. Se mostrarmos aos animais imagens de Drácula ou Frankenstein, eles riem e rasgam a figura. Mas o que faz os macaquinhos fugirem aterrorizados é a imagem de um macaco decapitado. Mais uma vez, é a perversão do comum que desperta o maior medo. (No capítulo 2, mencionamos que a teoria de espaço-tempo da consciência explica a natureza do humor, pois o cérebro simula o futuro de uma piada, e se surpreende ao ouvir o final inesperado. Isso explica também a natureza do terror.

O cérebro simula o futuro de um evento comum, trivial, e tem um choque quando de repente as coisas são horrivelmente pervertidas.)

Por esse motivo, os robôs continuarão a ter uma aparência meio infantil, mesmo que se aproximem da inteligência humana. Somente quando os robôs puderem agir realisticamente como humanos, seus projetistas poderão lhes dar uma aparência totalmente humana.

CONSCIÊNCIA DE SILÍCIO

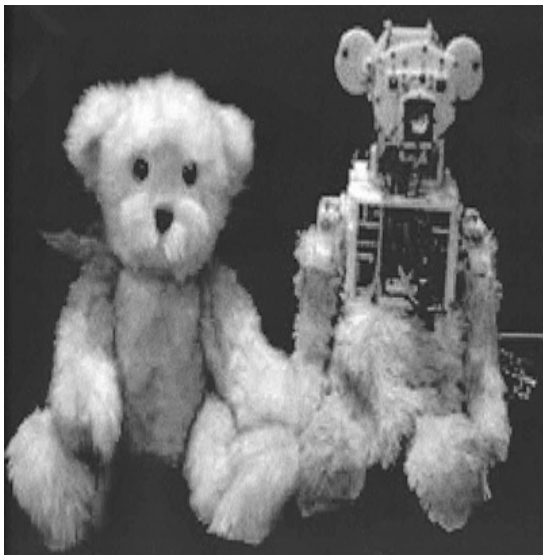
Como vimos, a consciência humana é uma colcha de retalhos imperfeita, confeccionada com as diversas capacidades desenvolvidas ao longo de milhões de anos de evolução. Se receberem informações sobre seu mundo físico e social, os robôs podem ser capazes de criar simulações semelhantes (ou, em alguns aspectos, até superiores) às nossas, mas uma consciência de silício pode diferir da nossa em duas áreas-chave: emoções e metas.

Historicamente, os pesquisadores da IA ignoraram o problema da emoção, considerando-a uma questão secundária. A meta era criar um robô lógico e racional, e não distraído e impulsivo. Consequentemente, a ficção científica dos anos 1950 e 1960 pôs em destaque robôs (e humanoides como Spock em *Jornada nas estrelas*) com cérebros perfeitos, totalmente lógicos.

Com o “vale estranho”, vimos que os robôs precisarão ter uma certa aparência para entrar em nossa casa, mas algumas pessoas argumentam que os robôs também devem ter emoções, para podermos criar laços, interagir produtivamente e também cuidar deles. Em outras palavras, os robôs precisarão do Nível II de consciência. Para conseguir isso, eles terão que reconhecer todo o espectro de emoções humanas. Ao analisar os movimentos faciais sutis de sobrancelhas, pálpebras, lábios, bochechas etc., um robô deverá ser capaz de identificar o estado emocional de um humano, principalmente o de seu dono. Uma instituição que se destacou na criação de robôs que reconhecem e imitam emoções é o MIT Media Laboratory. Tive o prazer de visitar o laboratório, perto de Boston, em várias ocasiões. É como visitar uma fábrica de brinquedos para gente grande. Para onde quer que a gente olhe, vemos equipamentos futuristas, de alta tecnologia, destinados a tornar nossa vida mais interessante, agradável e conveniente.

Olhando ao redor da sala, vi muitas imagens que entraram em filmes de Hollywood como *Minority Report – A nova lei* e *Inteligência Artificial*. Andando por aquele playground do futuro, encontrei dois robôs muito interessantes, Huggable e Nexi. Sua criadora, a dra. Cynthia Breazeal, me contou que eram robôs com metas específicas. Huggable (“abraçável”) é um ursinho-robô muito

fofo, feito para crianças. Ele consegue identificar emoções de crianças, tem câmeras de vídeo nos olhos, alto-falante na boca e sensores no pelo (assim, ele sabe se alguém está lhe fazendo cócegas, cutucando ou abraçando). Um dia, um robzinho desses pode ser um professor particular, babá, auxiliar de enfermagem, ou companheiro de brincadeiras.



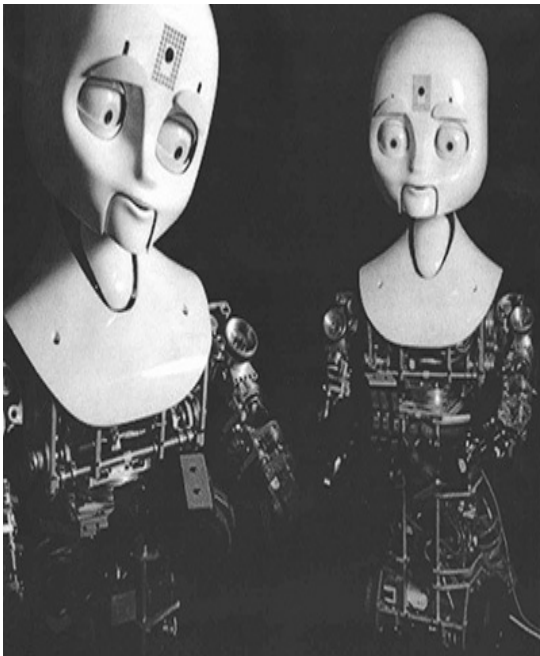


Figura 12. Huggable (topo) e Nexi (embaixo), dois robôs construídos no MIT Media Laboratory, projetados explicitamente para interagir com humanos através de emoções.

Nexi, por outro lado, é para adultos. Parece um pouco com o Pillsbury Doughboy. Tem uma carinha redonda, bochechuda, amigável, e olhos redondos que reviram para todo lado. Quando foi testado num asilo de idosos, os pacientes adoraram. Os velhinhos se afeiçoaram ao Nexi, o beijaram, conversaram com

ele, e sentiram saudades quando ele foi embora. (Ver Figura 12.)

A dra. Breazeal me disse que projetou Huggable e Nexi porque não estava satisfeita com seus robôs anteriores, que pareciam latas cheias de fios, motores e engrenagens. Para projetar um robô que pudesse interagir emocionalmente com as pessoas, a cientista precisou imaginar o que ele teria que fazer para se comportar como nós e criar laços afetivos. E ela não queria robôs presos no laboratório, mas que saíssem para o mundo real. O ex-diretor do MIT Media Laboratory, dr. Frank Moss, diz: “Foi por isso que em 2004 Breazeal decidiu que já era o momento de criar uma geração de robôs sociais, que pudessem viver em qualquer lugar: lares, escolas, hospitais, clínicas de idosos etc.”

Na Waseda University, no Japão, os cientistas estão trabalhando num robô com movimentos da parte superior do corpo representando emoções (medo, raiva, surpresa, alegria, nojo, tristeza) e que podem ouvir, cheirar, ver e tocar. Foram programados para atingir metas importantes, como se recarregar de energia e evitar situações perigosas. O objetivo é integrar os sentidos com as emoções, de modo que os robôs ajam de modo apropriado em diferentes situações.

Para não ficar atrás, a Comissão Europeia está financiando um projeto em andamento chamado Feelix Growing, para promover a Inteligência Artificial no Reino Unido, França, Suíça, Grécia e Dinamarca.

ROBÔS EMOTIVOS

Vou lhes apresentar o Nao.

Quando está feliz, ele estende os braços para cumprimentar com um grande abraço. Quando está triste, baixa a cabeça e parece desamparado, com os ombros curvados. Quando está com medo, se encolhe todo até que alguém lhe dê um tapinha tranquilizador na cabeça.

É como um menininho de 1 ano, só que robô. Nao tem 45 centímetros de altura, e parece muito com os robôs que vemos numa loja de brinquedos, como os Transformers, exceto que é um dos mais avançados robôs emotivos do mundo. Foi construído por cientistas da Universidade de Hertfordshire, no Reino Unido, em uma pesquisa financiada pela União Europeia.

É programado para externar emoções como felicidade, tristeza, medo, animação e orgulho. Enquanto outros robôs têm gestos faciais e verbais rudimentares para comunicar suas emoções, Nao conhece linguagem corporal, gestos e posturas. E até dança.

Ao contrário de outros robôs, especializados em apenas uma área emocional, Nao domina uma ampla escala de reações. Primeiro, Nao se concentra no rosto dos visitantes, identifica-os, e se lembra de suas interações anteriores com cada

um deles. Segundo, ele segue os movimentos deles. Por exemplo: ele segue o olhar deles, e sabe para onde estão olhando. Terceiro, ele começa a interagir e aprende a reagir aos gestos. Por exemplo: se alguém sorrir para ele ou der tapinhas em sua cabeça, ele sabe que é um sinal positivo. Como seu cérebro tem redes neurais, ele aprende a partir das interações com humanos. Quarto, Nao expressa emoções em resposta a suas interações com as pessoas (todas as suas respostas emocionais são pré-programadas, como num gravador de fita, mas ele decide qual é a emoção adequada à situação). Por fim, quanto mais Nao interage com um humano, mais ele aprende sobre os humores daquela pessoa, e mais fortes se tornam seus laços com ela.

Ele não tem só uma personalidade; pode ter várias. Já que Nao aprende com a interação com humanos, e cada interação é única, personalidades diferentes vão surgindo. Por exemplo: uma personalidade pode ser muito independente, dispensando orientação, e outra pode ser tímida, medrosa, apavorada com os objetos numa sala, exigindo constante intervenção humana.

A líder do projeto Nao é a dra. Lola Cañamero, cientista da computação na Universidade de Hertfordshire. Para dar início a esse ambicioso projeto, ela analisou a interação entre chimpanzés. Seu objetivo era reproduzir o mais fielmente possível o comportamento emocional de um chimpanzé de 1 ano de idade.

A dra. Cañamero vê aplicações imediatas para robôs emotivos. Assim como a dra. Breazeal, ela quer usá-los para reduzir a ansiedade de crianças internadas em hospitais. Ela diz: “Queremos explorar diversos papéis – os robôs podem ajudar as crianças a entender o tratamento, explicar o que elas devem fazer. Queremos ajudar as crianças a controlar a ansiedade.”

Outra possibilidade é que os robôs se tornem acompanhantes em asilos. Nao pode ser importante numa equipe médica. Algum dia esses robôs poderão ser companheiros de crianças e fazer parte da família.

“É difícil prever o futuro, mas não deve demorar muito para que seu computador seja um robô. Você poderá conversar, flertar, e até se irritar com ele – e ele entenderá você e suas emoções”, disse o dr. Terrence Sejnowski, do Salk Institute, perto de San Diego. Essa é a parte fácil. A parte difícil é calibrar as respostas do robô adequadas às informações. Se o dono está irritado ou descontente, o robô deve levar isso em conta ao responder.

EMOÇÕES: DETERMINANDO O QUE É IMPORTANTE

Além disso, os pesquisadores de IA começaram a perceber que as emoções podem ser uma chave para a consciência. Neurocientistas como o dr. Antonio

Damasio descobriram que, quando a conexão entre o lobo pré-frontal (que governa o pensamento racional) e os centros emocionais (por exemplo, o sistema límbico) é danificada, os pacientes não conseguem fazer julgamentos de valor. Eles ficam com a capacidade de decisão travada, sem conseguir resolver as mais simples (o que comprar, quando marcar um encontro, qual caneta usar), porque tudo tem o mesmo valor para eles. Portanto, vê-se que as emoções não são um luxo; são absolutamente essenciais e, sem elas, um robô terá dificuldade para determinar o que é importante e o que não é. As emoções, antes consideradas menos importantes para a Inteligência Artificial, estão agora assumindo uma importância central.

Se um robô se depara com um incêndio, ele pode salvar primeiro o computador, e não as pessoas, caso sua programação diga que documentos importantes não podem ser substituídos, e trabalhadores podem. É crucial que os robôs sejam programados para distinguir entre o que é importante e o que não é, e as emoções são atalhos que o cérebro pega para determinar isso rapidamente. Portanto, os robôs devem ser programados para ter um sistema de valores – que a vida humana é mais importante que objetos materiais, que crianças devem ser salvas primeiro numa emergência, que objetos mais caros têm mais valor que objetos baratos etc. Como os robôs não vêm equipados com valores, precisam ser carregados com uma lista imensa de julgamentos de valor.

O problema das emoções é que às vezes elas são irracionais, enquanto que os robôs são matematicamente precisos. Assim, a consciência de silício pode diferir da consciência humana em aspectos fundamentais. Por exemplo: os humanos têm pouco controle sobre as emoções, que surgem de forma repentina, pois se originam no sistema límbico, e não no córtex pré-frontal. Nossas emoções são quase sempre influenciadas. Vários testes mostram que tendemos a superestimar a capacidade de pessoas bonitas. Os belos tendem a ascender mais na sociedade e ter empregos melhores mesmo que sejam menos capazes que os outros, contrariando o ditado “Beleza não põe mesa”.

Da mesma forma, a consciência de silício pode não levar em conta sinais sutis que os humanos usam quando se encontram, como a linguagem corporal. Quando pessoas entram numa sala, os jovens geralmente mostram deferência aos mais velhos, e funcionários em posição inferior demonstram maior cortesia para com os superiores. Mostramos deferência no movimento do corpo, na escolha das palavras, nos gestos. A linguagem corporal é mais antiga que a própria linguagem falada, e é configurada de forma sutil no cérebro. Para interagir socialmente com humanos, os robôs terão que aprender esses sinais inconscientes.

Nossa consciência é influenciada por peculiaridades do nosso passado evolucionário, que os robôs não têm. Portanto, a consciência de silício pode não ter a mesma abertura e as mesmas peculiaridades que a nossa.

UM MENU DE EMOÇÕES

Já que a programação das emoções é feita externamente, e depois incorporada aos robôs, os fabricantes podem oferecer um menu de emoções escolhidas a dedo, conforme a necessidade, utilidade, e maior afinidade com o dono.

O mais provável é que os robôs sejam programados para ter apenas algumas emoções humanas, dependendo da situação. Talvez a emoção mais valorizada seja a lealdade. O dono quer que o robô execute suas ordens sem reclamar, que entenda, antecipe e satisfaça suas necessidades. A última coisa que o dono deseja é um robô cheio de atitudes, que dê respostas desaforadas, critique as pessoas e fique resmungando. Críticas construtivas são importantes, mas devem ser feitas com tato. E se o robô receber ordens conflitantes, deverá saber ignorar todas as que não sejam do seu dono.

A empatia é outra emoção valorizada pelo dono. Robôs com empatia entenderão os problemas de outros e se disponibilizarão para ajudar. Sabendo interpretar movimentos faciais e tom de voz, os robôs poderão identificar quando alguém está sofrendo e prestar a ajuda possível.

Estranhamente, o medo é outra emoção desejável. A evolução nos deu o sentimento de medo por uma razão, para evitarmos determinadas ameaças. Ainda que os robôs sejam feitos de aço, eles também devem temer certos perigos, como cair do alto de um edifício, ou entrar num incêndio. Um robô totalmente destemido, que se deixe destruir, será inútil.

Mas outras emoções, como a raiva, precisam ser excluídas, proibidas, ou fortemente reguladas. Se os robôs são construídos para ter grande força física, um robô furioso pode criar muitos problemas em lares e escritórios. A raiva pode impedi-lo de realizar suas tarefas e causar grandes danos a propriedades. (O propósito evolucionário da raiva é mostrar descontentamento. Isso pode ser feito de maneira racional, serena, sem fúria.)

Outra emoção que deve ser excluída é o desejo de estar no comando. Um robô mandão só criaria problemas, desafiando as ideias e desejos do dono (veremos mais adiante a importância dessa questão, quando discutirmos se um dia os robôs poderão dominar os humanos). Portanto, o robô terá que se curvar aos desejos do dono, mesmo que não sejam os mais aconselháveis.

Talvez a emoção mais difícil de programar seja o humor, que é capaz de unir dois completos estranhos. Uma simples piada pode desarmar uma situação tensa ou acentuá-la ainda mais. O mecanismo do humor é simples: o final da anedota é inesperado. Mas as sutilezas do humor são enormes. Muitas vezes avaliamos o outro pelo modo como reage a certas piadas. Se os humanos usam piadas como medida para avaliar outros humanos, seria importante criar um robô que possa achar uma piada engraçada. O presidente Ronald Reagan, por exemplo, era famoso por descontrair com uma tirada as situações mais difíceis. Ele tinha uma

vasta coleção de piadas, tiradas e comentários afiados porque entendia o poder do humor. (Alguns comentaristas políticos acreditam que ele venceu o debate com Walter Mondale na campanha para a presidência quando foi perguntado se não era velho demais para ser presidente, e Reagan respondeu que não usaria a juventude do adversário contra ele.) O riso inapropriado pode ter consequências desastrosas (e pode até ser sinal de doença mental). O robô precisará discernir entre rir *com* alguém ou *de* alguém. (Os atores conhecem bem a natureza diversificada do riso. Sabem criar risadas que representam horror, cinismo, alegria, raiva, tristeza etc.) Portanto, pelo menos até que a teoria da Inteligência Artificial esteja mais desenvolvida, os robôs devem ficar longe do humor e das risadas.

PROGRAMANDO EMOÇÕES

Até agora evitamos discutir a difícil questão sobre como essas emoções seriam programadas num computador. Devido à complexidade, provavelmente terão que ser programadas em estágios.

Primeiro, o mais fácil é identificar uma emoção analisando as expressões do rosto, lábios, sobrancelhas e tom de voz de alguém. Atualmente, a tecnologia de reconhecimento facial já é capaz de criar um dicionário de emoções, atribuindo determinados significados a certas expressões faciais. Na verdade, esse processo remonta a Charles Darwin, que passou bastante tempo catalogando emoções comuns a animais e humanos.

Segundo, o robô deve responder rapidamente a cada emoção. Isso também é fácil. Se alguém rir, o robô sorri. Se alguém estiver com raiva, o robô sai de perto para evitar conflito. O robô deverá ter uma vasta enciclopédia de emoções programadas e encontrar uma resposta rápida para cada uma.

O terceiro estágio é o mais complexo porque envolve determinar as motivações por trás da emoção aparente. Isso é difícil, pois há uma série de situações que podem deflagrar uma determinada emoção. Uma risada pode significar que alguém está feliz, escutou uma piada, ou viu uma pessoa cair. Ou pode ser que esteja nervoso, ansioso, ou insultando alguém. Da mesma forma, se alguém grita, pode significar uma emergência, ou pode ser apenas uma reação de alegria e surpresa. Determinar a razão por trás de uma emoção exige uma habilidade difícil até para os humanos. Para conseguir isso, o robô precisará ter uma lista de várias razões possíveis para uma emoção, e determinar a que faz mais sentido. Isso significa tentar encontrar a razão por trás da emoção que se encaixe melhor em sua base de dados.

Quarto, uma vez determinada a origem da emoção, o robô terá que dar a

resposta adequada. Isso também é difícil, pois há várias respostas possíveis, e a resposta errada pode piorar a situação. O robô já tem em sua programação uma lista das reações possíveis à emoção original. Ele precisará calcular qual é a mais apropriada, o que significa simular o futuro.

OS ROBÔS VÃO MENTIR?

Normalmente, pensamos em robôs friamente analíticos e racionais, sempre dizendo a verdade. Mas, uma vez integrados à sociedade, eles provavelmente aprenderão a mentir ou, pelo menos, a ter tato suficiente para refrear seus comentários.

Em nossa vida, várias vezes por dia nos deparamos com situações em que precisamos dizer uma mentira branca. Se alguém nos pergunta se está com boa aparência, nem sempre ousamos dizer a verdade. A mentira branca é como um lubrificante para a sociedade funcionar bem. Se formos subitamente obrigados a dizer toda a verdade (como Jim Carrey em *O mentiroso*), acabaremos criando muita confusão e magoando os outros. Muitos ficarão ofendidos se você disser a verdade sobre a aparência deles, ou o que você realmente acha. Você vai ser demitido pelo chefe. Abandonado pela pessoa amada. Evitado pelos amigos. Estranhos lhe darão uma bofetada. É melhor guardar para si alguns pensamentos.

Da mesma forma, os robôs precisarão aprender a mentir ou esconder a verdade, ou acabarão ofendendo as pessoas e dispensados pelos donos. Se um robô disser a verdade numa festa, pode ficar mal para o dono e criar uma confusão. Portanto, se alguém pedir a opinião do robô, ele tem que saber ser evasivo, agir com diplomacia, e responder com tato. Ele pode se esquivar da questão, mudar de assunto, responder com um clichê, com outra pergunta ou uma mentira branca (tudo o que os *chat-bots* têm aperfeiçoado). Isso significa que o robô já estará programado com uma lista de respostas evasivas possíveis, e deverá escolher a que criar menos complicações.

Uma das poucas vezes em que um robô deve dizer toda a verdade é ao responder a uma pergunta do dono, que precisa estar preparado para uma resposta franca. Talvez outra ocasião em que o robô deva ser sincero seja numa investigação policial, quando a verdade é absolutamente necessária. A não ser nessas circunstâncias, os robôs devem ser capazes de mentir livremente ou de omitir a verdade para manter as engrenagens da sociedade em funcionamento.

Em outras palavras, os robôs têm que ser socializados tal como os adolescentes.

OS ROBÔS VÃO SENTIR DOR?

Os robôs em geral são programados para tarefas cansativas, sujas e perigosas. Não há motivo para deixarem de fazer trabalhos repetitivos ou sujos, pois não são programados para sentir tédio ou nojo. O problema surge quando o robô tem que enfrentar situações perigosas. Nesse ponto, talvez seja preciso realmente programá-los para sentir dor.

Desenvolvemos a sensação da dor porque nos ajudou a sobreviver em ambientes perigosos. Há pessoas que nascem sem a capacidade de sentir dor, devido a um defeito genético chamado analgesia congênita. À primeira vista, pode parecer uma sorte, pois essas crianças não choram quando se machucam, mas na verdade é uma maldição. As crianças com esse defeito genético têm sérios problemas, como morder e cortar fora partes da língua, sofrer queimaduras graves e cortes, que podem levar à amputação dos dedos. A dor é um alerta de perigo, avisando para tirar a mão do fogo aceso ou parar de correr com o tornozelo torcido.

Em algum momento, os robôs terão que ser programados para sentir dor, ou não saberão evitar situações perigosas. A primeira sensação de dor deverá ser a fome (isto é, a sensação de necessidade de energia elétrica). Quando as baterias estiverem fracas, eles vão ficar desesperados, sabendo que seus circuitos logo se desligarão, deixando o trabalho incompleto. Quanto menos carga tiver a bateria, mais ansiosos eles ficarão.

Além disso, por mais fortes que sejam, os robôs podem pegar acidentalmente um objeto pesado demais e quebrar seus membros. Podem sofrer superaquecimento ao trabalhar com metal fundido numa metalúrgica, ou ao entrar num prédio incendiado para auxiliar os bombeiros. Nesses casos, os sensores de temperatura e tensão precisam dar o alerta de que estão excedendo os limites de suas especificações.

Mas quando a sensação de dor é incluída no menu de emoções, levanta imediatamente questões éticas. Muita gente acredita que não devemos infligir dor desnecessária aos animais e tem a mesma opinião a respeito dos robôs. Isso abre espaço para os direitos dos robôs. Haverá leis para limitar a exposição de robôs à dor e ao perigo. As pessoas não se importam se os robôs estão executando tarefas entediantes ou sujas, mas, se eles sentirem dor ao executar uma tarefa perigosa, pode ter início um lobby para leis de proteção aos robôs. Isso pode levar a um conflito legal entre donos e fabricantes exigindo um aumento do nível da dor que os robôs poderão tolerar, e especialistas em ética querendo o contrário.

Isso, por sua vez, pode deflagrar outros debates sobre os direitos dos robôs. Os robôs poderão ter propriedades? O que acontecerá se eles ferirem alguém acidentalmente? Poderão ser processados ou punidos? Quem será o responsável num processo legal? Um robô poderá ser dono de outro robô? Essa discussão

levanta outra questão espinhosa: Os robôs devem ser programados com um senso de ética?

ROBÔS ÉTICOS

A princípio, a ideia de robôs éticos parece perda de tempo e esforço. Entretanto, essa questão adquire outro tom quando consideramos que os robôs poderão tomar decisões de vida ou morte. Como serão fisicamente fortes e capazes de salvar vidas, terão que fazer escolhas éticas instantâneas a respeito de quem salvar primeiro.

Digamos que haja um terremoto catastrófico, e crianças estão presas nos escombros de um prédio. Onde o robô deverá empregar sua energia? Deve tentar salvar o maior número de crianças? Ou as crianças menores? Ou as mais vulneráveis? Se os escombros forem pesados demais, o robô pode danificar sua constituição eletrônica. Nesse caso, há mais um impasse ético: como calcular o número de crianças que salvará versus os danos que sua constituição eletrônica aguentará?

Sem uma programação adequada, o robô pode simplesmente parar, esperando que um humano tome a decisão, “e assim perderá um tempo precioso”. Portanto, alguém tem que programá-lo antevendo situações em que ele precise tomar automaticamente a decisão “certa”.

As decisões éticas terão que ser pré-programadas no computador logo no início, pois não existe lei matemática que atribua um valor ao salvamento de um grupo de crianças. Nessa programação deverá constar uma longa lista de coisas e situações classificadas por ordem de importância. É um trabalho muito cansativo. Às vezes um ser humano leva a vida inteira para aprender uma lição ética, mas um robô deve aprendê-las antes de sair da fábrica para poder se integrar à sociedade. Só as pessoas conseguem fazer isso, e mesmo assim os dilemas éticos às vezes nos confundem. E isso levanta questões: Quem tomará as decisões? Quem decide a ordem em que os robôs salvarão vidas humanas?

A questão da tomada de decisões provavelmente será resolvida por uma combinação de leis e mercado. Será preciso haver leis para, no mínimo, determinar uma ordem de prioridades de salvamento numa emergência. Mas, além dessa, há milhares de questões éticas menores. As decisões sobre questões mais sutis deverão ser decididas pelo mercado e o senso comum.

Se você trabalha para uma empresa de segurança, com clientes de alta posição, terá que ensinar o robô a salvar as pessoas numa ordem exata em diferentes situações, com base em considerações como cumprimento do dever, mas também orçamento.

O que acontece se um bandido comprar um robô para cometer um crime? Isso levanta outra questão: o robô deverá descumprir a ordem do seu dono para infringir a lei? Nos exemplos anteriores, vimos que os robôs serão programados para compreender as leis e tomar decisões éticas. Portanto, se ele julgar que o dono está lhe pedindo para infringir a lei, deve estar liberado para desobedecer.

Há também o dilema ético colocado por robôs que refletem as crenças do dono, que podem divergir das normas morais e sociais. As “guerras culturais” que vemos na sociedade de hoje só tendem a se intensificar se tivermos robôs que sigam as crenças e opiniões do dono. Em certo sentido, esse conflito é inevitável. Os robôs são extensões dos sonhos e desejos de seus criadores e, quando forem sofisticados o suficiente para tomar decisões morais, eles tomarão.

As diferenças na sociedade poderão ser ressaltadas quando os robôs começarem a ter comportamentos que desafiam nossos valores e objetivos. Robôs de jovens que gostam de bandas barulhentas podem entrar em conflito com robôs de idosos residentes num bairro sossegado. O primeiro grupo de robôs pode ser programado para amplificar o som, e o segundo grupo ser programado para reduzir os ruídos ao mínimo possível. Robôs de religiosos fundamentalistas podem ter desavenças com robôs de ateus. Robôs de diferentes nações e culturas serão projetados para refletir os costumes de sua sociedade, o que pode gerar choques (se acontece entre humanos, imagine entre robôs). Então, como será possível programar robôs eliminando esses conflitos? Não será. Os robôs simplesmente refletirão as tendências e preconceitos dos criadores. No fim das contas, as diferenças culturais e éticas entre robôs serão resolvidas na justiça. Não existe lei da física ou da ciência que determine essas questões morais e, portanto, será preciso criar leis para resolver os conflitos sociais. Os robôs não podem resolver dilemas morais criados pelos humanos. Na verdade, só podem intensificá-los.

Mas, se os robôs puderem tomar decisões éticas e legais, também poderão ter e entender sensações? Se conseguirem salvar alguém, sentirão alegria? Terão preferências, como gostar da cor vermelha? Analisar friamente a ética de quem salvar é uma coisa, mas entender e sentir é outra. Os robôs poderão sentir?

OS ROBÔS PODERÃO ENTENDER E SENTIR?

Ao longo dos séculos, muitas teorias questionaram se uma máquina poderia pensar e sentir. Minha filosofia própria se chama “construtivismo”, isto é, em vez de debater a questão indefinidamente, o que não leva a nada, devemos empregar nossa energia na criação de um autômato para ver até onde chegamos. De outro modo, vamos acabar em intermináveis debates filosóficos, sem chegar a uma

conclusão. A vantagem da ciência é que, quando tudo já foi dito e feito, podemos realizar experimentos para decidir a questão.

Assim, para decidir essa questão, se um robô pode pensar ou não, a solução está em construir um. No entanto, há quem defenda que máquinas jamais serão capazes de pensar como os humanos. O argumento mais forte é que, apesar de conseguir manipular fatos mais rapidamente que os humanos, um robô não “entende” o que está manipulando. Embora seja capaz de processar sentidos (cor, som) melhor que um humano, não consegue realmente “sentir” ou “experimentar” a essência desses sentidos.

Por exemplo: o filósofo David Chalmers dividiu os problemas da IA em duas categorias, Problemas Fáceis e Problemas Difíceis. Para ele, Problemas Fáceis são criar máquinas que imitam cada vez mais as habilidades humanas, como jogar xadrez, somar números, reconhecer certos padrões etc. Os Problemas Difíceis envolvem criar máquinas que entendam sentimentos e sensações subjetivas, que são chamados qualia.

De acordo com essa abordagem, do mesmo modo que é impossível ensinar a um cego o que é a cor vermelha, um robô jamais será capaz de vivenciar a sensação subjetiva da cor vermelha. E um computador pode traduzir palavras do chinês para o inglês com grande fluência, mas jamais entenderá o que está traduzindo. Nesse cenário, os robôs são excelentes gravadores ou máquinas de somar, capazes de recitar e manipular informação com uma precisão incrível, mas sem entender nada.

Esses argumentos devem ser levados a sério, mas há outra maneira de ver a questão dos qualia e da experiência subjetiva. No futuro, é provável que haja uma máquina capaz de processar uma sensação, como a cor vermelha, muito melhor que os humanos. Será capaz de descrever as propriedades físicas do vermelho, e até usá-lo num verso poético melhor que um humano. Mas o robô “sente” a cor vermelha? Isso é irrelevante, pois a palavra “sentir” não é bem definida. Algum dia, a descrição da cor vermelha feita por um robô poderá ser melhor que a humana, e o robô poderá perguntar: Os humanos entendem realmente a cor vermelha? Talvez os humanos não sejam capazes de entender a cor vermelha com todas as nuances e sutilezas que um robô consegue captar.

Como disse o behaviorista B. F. Skinner, “o problema real não é se as máquinas pensam, mas se os homens pensam”.

Da mesma forma, é uma questão de tempo até que um robô seja capaz de definir palavras chinesas e usá-las num contexto de modo muito melhor que qualquer humano. Nesse ponto, é irrelevante se o robô “entende” ou não chinês. Na prática, o computador conhecerá a língua chinesa melhor que qualquer humano. Em outras palavras, a palavra “entender” não tem uma definição exata.

Um dia, quando os robôs tiverem superado nossa capacidade de manipular essas palavras e sensações, será irrelevante se o robô as “sente” ou “entende”. A

questão deixará de ter importância.

Como disse o matemático John von Neumann: “Na matemática, não entendemos as coisas. Só nos acostumamos a elas.”

Assim, o problema não está na máquina, mas na natureza da linguagem humana, nas palavras que não possuem uma definição exata e têm significados diferentes para pessoas diferentes. Certa vez perguntaram ao grande físico quântico Niels Bohr como alguém poderia entender os profundos paradoxos da teoria quântica. A resposta, ele disse, está em como você define “entender”.

O dr. Daniel Dennett, filósofo na Tufts University, escreveu: “Não é possível haver um teste objetivo para diferenciar um robô esperto de uma pessoa consciente. Então, você escolhe: ou permanece fixado no Problema Difícil, ou desiste e esquece o assunto. Deixa pra lá.”

Em outras palavras, não existe o Problema Difícil.

Para a filosofia construtivista, o que interessa não é discutir se uma máquina pode “sentir” a cor vermelha, mas construir a máquina. Nesse cenário, há um continuum de níveis descrevendo as palavras “entender” e “sentir” (isso significa que pode até ser possível atribuir valores numéricos a graus de entendimento e sensação). Num polo, temos os robôs atuais, desajeitados, que conseguem manipular alguns símbolos, mas não vão muito além disso. No polo oposto, temos os humanos, que se orgulham de sentir os qualia. Mas, com o tempo, os robôs serão capazes de descrever sensações melhor que nós, em todos os níveis. Então será óbvio que os robôs entendem.

Essa era a filosofia por trás do famoso teste de Turing, de Alan Turing. Ele previu que algum dia haveria uma máquina capaz de responder como um humano a qualquer pergunta: “Um computador vai merecer ser chamado de inteligente quando puder fazer um humano acreditar que ele é humano.”

A definição de Francis Crick, físico laureado pelo Prêmio Nobel, é melhor ainda. Ele observou que no século passado os biólogos tiveram discussões acaloradas sobre “O que é a vida?”. Agora, com o entendimento que adquirimos sobre o DNA, os cientistas percebem que a questão não é bem definida. Essa simples questão tem muitas variações, camadas e complexidades. A pergunta “O que é a vida?” simplesmente desapareceu. O mesmo pode vir a acontecer com o entendimento e a sensação.

ROBÔS CONSCIENTES DE SI

Que passos devem ser dados para que computadores como o Watson tenham consciência de si? Para responder a essa pergunta, vamos voltar à definição de consciência de si: é a capacidade de se colocar em um modelo do ambiente e

fazer simulações desse modelo no futuro para atingir um objetivo. Esse primeiro passo exige um nível muito alto de senso comum, a fim de prever vários eventos. Portanto, o robô deve se colocar nesse modelo, o que requer um entendimento das várias possibilidades de ação que ele poderá escolher.

Na Universidade Meiji, os cientistas deram o primeiro passo para a criação de um robô com consciência de si. É um projeto ambicioso, mas eles acham que podem executá-lo criando robôs com a Teoria da Mente. Começaram com a construção de dois robôs. O primeiro foi programado para executar certos movimentos. O segundo foi programado para observar o primeiro e imitá-lo. Conseguiram criar um segundo robô que imitava sistematicamente o comportamento do primeiro, somente observando-o. É a primeira vez na história que se constrói um robô especificamente para ter consciência de si. O segundo robô tem a Teoria da Mente, isto é, é capaz de observar outro robô e imitar seus movimentos.

Em 2012, o segundo passo foi dado por cientistas da Universidade de Yale, ao criarem um robô que passou no teste do espelho. Quando animais são colocados diante de um espelho, a maioria acha que a imagem no espelho é de outro animal. Como vimos, apenas alguns animais passaram no teste do espelho, percebendo que a imagem era um reflexo deles mesmos. Os cientistas de Yale criaram um robô chamado Nico, que parece um esqueleto desengonçado, feito de arames retorcidos, braços mecânicos e dois olhos salientes no alto. Quando foi colocado na frente do espelho, Nico não só reconheceu sua imagem, mas conseguiu deduzir a localização de objetos na sala ao ver a imagem deles no espelho. Isso é similar ao que fazemos ao olhar pelo espelho retrovisor e inferir a localização dos objetos lá atrás.

O programador de Nico, Justin Hart, disse: “Até onde sabemos, é o primeiro sistema robótico a usar o espelho dessa maneira, e representa um passo significativo na direção de uma arquitetura coesa, que permite ao robô aprender sobre seu corpo e aparência por meio da auto-observação, e é uma aptidão importante, exigida para passar no teste do espelho.”

Como os robôs da Universidade Meiji e da Universidade de Yale representam a mais avançada tecnologia na construção de robôs com consciência de si, é fácil perceber que os cientistas têm um longo caminho até criarem um robô com consciência de si semelhante à humana.

A pesquisa está apenas começando, porque nossa definição de consciência de si requer que o robô utilize essa informação para criar simulações do futuro, o que está muito além da capacidade de Nico, ou de qualquer outro robô.

Isso levanta uma questão importante: Como um computador pode adquirir total consciência de si? Na ficção científica, costumamos nos deparar com uma situação em que a internet de repente se torna consciente de si, como no filme *O exterminador do futuro*. Dado que a internet está conectada a toda a infraestrutura

da sociedade moderna (por exemplo, sistema de esgoto, eletricidade, telecomunicações, armamentos), uma internet consciente de si poderia assumir facilmente o controle da sociedade. Ficariamos impotentes diante dessa situação. Alguns cientistas opinam que isso pode acontecer como exemplo de um “fenômeno emergente”, ou seja, quando se reúne um número grande de computadores, pode ocorrer uma súbita transição de fase para um estágio superior, sem qualquer fornecimento de dados do exterior.

Isso diz tudo e não diz nada, porque deixa de fora todos os importantes passos intermediários. É como falar que uma via expressa pode subitamente ter consciência de si se tiver muitas pistas.

Neste livro, já que demos uma definição de consciência e de consciência de si, deve ser possível listar os passos para que a internet se torne consciente de si.

Primeiro, uma internet inteligente tem que fazer continuamente modelos de seu lugar no mundo. Em princípio, essa informação pode ser programada de forma externa. Isso inclui descrever o mundo externo, ou seja, a Terra, as cidades, os computadores, e tudo isso se pode encontrar na própria internet.

Segundo, a internet tem que se colocar nesse modelo. Essa informação também é fácil de conseguir. Inclui todas as especificações da internet (o número de computadores, nós, linhas de transmissão etc.) e sua relação com o mundo externo.

Mas o terceiro passo é muito mais difícil. Significa fazer simulações contínuas desse modelo no futuro, de acordo com um objetivo. É aí que esbarramos numa barreira. A internet não é capaz de fazer simulações do futuro, e não tem objetivos. Mesmo no mundo científico, simulações do futuro são geralmente feitas com poucos parâmetros (por exemplo, simulando a colisão de dois buracos negros). Fazer uma simulação de um modelo do mundo contendo a internet está muito além das possibilidades de programação disponível hoje. Seria preciso incorporar todas as leis do senso comum, todas as leis da física, química e biologia, bem como fatos sobre o comportamento humano e a sociedade humana.

Além disso, essa internet inteligente precisa ter um objetivo. Hoje, é apenas uma via expressa passiva, sem qualquer direção ou propósito. Claro que, em princípio, pode-se impor um objetivo à internet. Mas consideremos o seguinte problema: Podemos criar uma internet cujo objetivo seja a autopreservação?

Esse seria o objetivo mais simples de todos, mas ninguém sabe fazer nem mesmo essa programação. Um programa desses, por exemplo, impossibilita o desligamento da internet quando tiramos a tomada da parede. No momento, a internet é totalmente incapaz de reconhecer uma ameaça à sua existência, quanto mais uma conspiração para detê-la.

Por exemplo: uma internet capaz de detectar ameaças a sua existência precisa ser capaz de identificar tentativas de corte de eletricidade, cortes de linhas de

comunicação, destruição dos servidores, tentativas de desabilitar suas conexões por fibras óticas e por satélite etc. Além disso, uma internet capaz de se defender contra esses ataques precisa saber adotar contramedidas para cada situação, simulando atentados no futuro. Nenhum computador do mundo é capaz de fazer sequer uma fração disso.

Em outras palavras, talvez seja possível criar robôs conscientes de si, e até uma internet consciente de si, mas esse dia está num futuro longínquo, quem sabe, no final deste século?

Mas suponhamos que esse dia tenha chegado, e que encontramos robôs conscientes de si por aí. Se um desses robôs tem objetivos compatíveis com os nossos, esse tipo de Inteligência Artificial não traz nenhum problema. Mas o que acontecerá se os objetivos forem diferentes? O medo é de que os humanos sejam dominados e escravizados pelos robôs conscientes de si. Se tiverem uma capacidade superior de simular o futuro, os robôs poderão tramar os resultados de várias situações de modo a encontrar o melhor meio de subjugar a humanidade.

Uma forma de conter o problema é assegurar que os objetivos dos robôs sejam benevolentes. Como vimos, não basta simular o futuro. As simulações devem servir a um objetivo. Se o objetivo do robô é sua mera preservação, ele vai tentar se defender quando quisermos desligá-lo, e isso vai ser complicado para a humanidade.

OS ROBÔS ASSUMIRÃO O CONTROLE?

Na maioria dos contos de ficção científica, os robôs se tornam perigosos devido ao desejo de assumir o controle. A palavra “robô” vem do tcheco e significa “trabalhador”, tendo surgido pela primeira vez em 1920, na peça *R.U.R., O nascimento do robô* [*Rosumovi Univerzální Roboti*], de Karel Čapek, em que cientistas criam uma raça de seres mecânicos de aparência idêntica à dos humanos. Milhares desses robôs executam os trabalhos mais servis e perigosos, porém, são muito maltratados e um dia se rebelam e destroem a raça humana. Apesar de terem dominado o mundo, esses robôs têm uma falha: não conseguem se reproduzir. Mas no final da peça, dois robôs se apaixonam. Assim, talvez, venha a surgir um novo ramo da “humanidade”.

Uma trama mais realista é apresentada no filme *O exterminador do futuro*, em que os militares criam uma rede de supercomputadores, chamada Skynet, que controla todo o estoque de suprimento nuclear dos Estados Unidos. Certo dia, a Skynet se torna consciente. Os militares tentam desligá-la, mas então se dão conta de uma falha grave na programação: a rede foi projetada para se proteger,

e a única maneira de fazer isso é eliminar o problema: a humanidade. A Sky net desencadeia uma guerra nuclear, reduzindo a humanidade a um bando de desajustados e rebeldes lutando contra máquinas exterminadoras.

Certamente, é possível que os robôs se tornem uma ameaça. Hoje em dia, o drone (veículo aéreo não tripulado) Predator é capaz de alvejar vítimas com exatidão, mas é controlado por um *joystick* a milhares de quilômetros de distância. Segundo o *New York Times*, as ordens de ataque são dadas diretamente pelo presidente dos Estados Unidos. Mas, no futuro, o Predator poderá ter uma tecnologia de reconhecimento facial que lhe permita atirar quando tiver 99% de certeza da identidade do alvo. Sem intervenção humana, ele poderá usar automaticamente essa tecnologia para atacar qualquer um que se encaixe no perfil.

Pois bem, suponhamos que um drone desses tenha um colapso que cause alterações nos programas de reconhecimento facial. Ele pode se tornar um robô vilão com permissão para matar qualquer um que cruze seu caminho. Pior ainda, imagine uma esquadrilha desses drones controlada por um comando central. Se houver uma pane num único transistor do computador central, a esquadrilha inteira pode sair matando a esmo.

Um problema mais sutil seria quando robôs se comportam perfeitamente bem, sem panes, mas com uma pequenissima falha fatal na programação e nos objetivos. Para um robô, a autopreservação é um objetivo importante. Mas ser útil aos humanos, também. O problema surge quando esses dois objetivos entram em contradição.

No livro *Eu, Robô*, o sistema do computador decide que os humanos são autodestrutivos, com suas infundáveis guerras e atrocidades, e que a única maneira de proteger a raça humana é assumir o controle e criar uma benevolente ditadura da máquina. A contradição aqui não está entre dois objetivos, mas num único objetivo não realista. Esses robôs assassinos não têm um defeito – eles chegam à conclusão lógica de que a única maneira de preservar a humanidade é assumir o comando.

Uma solução para esse problema é criar uma hierarquia de objetivos. Por exemplo: o desejo de ajudar os humanos deve vir antes da autopreservação. Esse tema foi explorado no filme *2001: Uma odisseia no espaço*, em que HAL 9000 é um computador consciente, capaz de conversar tranquilamente com os humanos. Mas as ordens dadas a HAL são autocontraditórias, e a lógica impede que sejam executadas. Ao tentar executar o impossível, HAL passa dos limites, enlouquece, e a única solução para obedecer aos comandos contraditórios dos humanos imperfeitos é eliminá-los.

A melhor solução talvez seja criar uma lei da robótica, determinando que os robôs não podem fazer mal à raça humana mesmo que haja contradições em suas diretivas. Eles devem ser programados para ignorar contradições menores

em suas ordens e preservar sempre a lei suprema. Mas isso, na melhor das hipóteses, ainda pode ser um sistema imperfeito. Por exemplo: se o objetivo central do robô for proteger a humanidade, com prioridade sobre quaisquer outros objetivos, tudo vai depender de como o robô entende a palavra “proteger”. A definição mecânica dessa palavra pode ser diferente da nossa.

Alguns, como o dr. Douglas Hofstadter, cientista cognitivo da Universidade de Indiana, não temem a ameaça dos robôs. Quando o entrevistei, ele disse que os robôs são nossos filhos, e, então, por que não amá-los como amamos nossos próprios filhos? Nossa atitude, disse ele, é amar os filhos, mesmo sabendo que eles nos controlarão.

Quando entrevistei o dr. Hans Moravec, ex-diretor do Laboratório de AI da Carnegie Mellon University, ele concordou com Hofstadter. Em seu livro *Homens e robôs*, ele diz: “Livres do lento caminhar da evolução biológica, os filhos de nossa mente terão a liberdade de crescer para enfrentar os desafios imensos e fundamentais no universo maior. (...) Nós, humanos, vamos nos beneficiar do trabalho deles por algum tempo, mas, (...) assim como o nossos filhos naturais, eles buscarão a própria sorte, enquanto nós, os velhos pais, vamos desvanecendo silenciosamente.”

Outros, pelo contrário, acham essa solução horrível. Talvez o problema possa ser resolvido se modificarmos nossos objetivos e prioridades antes que seja tarde demais. Se os robôs são nossos filhos, temos que “ensiná-los” a ser benevolentes.

INTELIGÊNCIA ARTIFICIAL AMIGÁVEL

Como robôs são criaturas mecânicas feitas em laboratório, a natureza assassina ou amigável de tal máquina depende da direção tomada pelas pesquisas em IA. Muitos financiamentos vêm das Forças Armadas, que têm a incumbência específica de vencer guerras, portanto, robôs assassinos são uma possibilidade real.

Contudo, já que 30% dos robôs comerciais são fabricados no Japão,^[4] outra possibilidade é que sejam projetados para serem parceiros e trabalhadores. Esse objetivo é viável se o setor do consumo controlar a pesquisa robótica. Na filosofia de “IA amigável”, os inventores criarão robôs programados desde o início para serem benéficos para os humanos.

Culturalmente, a abordagem japonesa é diferente da ocidental. Enquanto as crianças ocidentais tremem de medo de robôs do tipo Exterminador, as japonesas são influenciadas pelo xintoísmo, pregando que tudo é dotado de espírito, inclusive robôs. Em vez de sentir estranheza diante de robôs, as crianças japonesas dão gritinhos de prazer ao encontrar um deles. Não é de admirar,

portanto, a proliferação de robôs no mercado e nas residências do Japão. Os robôs cumprimentam os clientes nas grandes lojas e dão aulas na TV. Há até uma peça de teatro no Japão estrelada por um robô.

O Japão tem outro motivo para acolher os robôs. Eles serão os futuros enfermeiros de um país idoso, onde 21% da população têm acima de 65 anos. O Japão está envelhecendo mais rapidamente que qualquer outra nação. Em certo sentido, o país é como um desastre ferroviário em câmera lenta. Há três fatores demográficos em ação. Primeiro, as mulheres japonesas têm uma expectativa de vida mais longa que qualquer outro grupo étnico no mundo. Segundo, o Japão tem a menor taxa de natalidade do mundo. Terceiro, a política de imigração é muito restritiva, e 99% da população é de japoneses legítimos. Não dispondo de imigrantes jovens, o país irá depender de robôs para cuidar dos velhos.

Esse problema não se restringe ao Japão; a Europa vem a seguir. Itália, Alemanha, Suíça e outros países europeus sentem uma pressão demográfica semelhante. As populações do Japão e da Europa poderão sofrer uma grave redução em meados do século. Os Estados Unidos não ficam muito atrás. A taxa de natalidade de cidadãos nascidos no país caiu drasticamente nas últimas décadas, mas a imigração manterá o crescimento norte-americano neste século. Em outras palavras, é uma aposta de trilhões de dólares para ver se os robôs poderão nos salvar desses pesadelos demográficos.

O Japão é líder mundial na criação de robôs capazes de entrar em nossa vida cotidiana. Os japoneses criaram robôs que cozinham (alguns conseguem preparar uma tigela de macarrão em um minuto e quarenta segundos). Em um restaurante, é possível fazer o pedido num tablet e o robô cozinheiro imediatamente entra em ação. Ele tem em dois grandes braços mecânicos, que pegam tigelas, colheres, facas, e preparam a refeição. Alguns robôs-cozinheiros têm até aparência humana.

Há também robôs musicais. Um deles tem “pulmões de acordeão”, que lhe permitem fazer música bombeando ar através de um instrumento. Há robôs-empregadas domésticas. É só separar a roupa lavada que o robô dobra. Existe até um robô que fala, pois tem pulmões artificiais, lábios, língua e cavidade nasal. A Sony Corporation, por exemplo, construiu o robô AIBO, que parece um cachorro e registra várias emoções quando é acariciado. Alguns futurólogos preveem que a indústria robótica pode vir a ser tão grande quanto a automobilística é hoje.

O mais importante aqui é que os robôs não são necessariamente programados para destruir e dominar. O futuro da Inteligência Artificial depende de nós.

Mas alguns críticos da IA amigável alegam que os robôs poderão tomar o controle, não que sejam agressivos, e sim por erro nosso ao criá-los. Em outras palavras, se um robô vier a dominar, será porque foi programado com objetivos conflitantes.

“EU SOU UMA MÁQUINA”

Quando entrevistei o dr. Rodney Brooks, ex-diretor do Laboratório de Inteligência Artificial do MIT e cofundador do iRobot, perguntei se ele achava que algum dia as máquinas iriam nos dominar. Ele me disse que precisamos aceitar que todos nós somos máquinas. Isso significa que algum dia seremos capazes de construir máquinas tão vivas quanto nós. Mas ele alerta que teremos que abrir mão do conceito de que somos “especiais”.

Essa evolução da perspectiva humana teve início com Nicolau Copérnico, quando descobriu que a Terra não era o centro do universo, e que girava em torno do Sol. Continuou com Darwin, ao mostrar que somos semelhantes aos animais em termos de evolução. E continuará no futuro – prosseguiu o dr. Brooks – quando descobrirmos que somos máquinas, apesar de feitos de carne e osso, e não de peças.

Ele acredita que será uma mudança importantíssima em nossa visão de mundo aceitar que nós também somos máquinas. “Não gostamos de abrir mão de nossa condição de especiais, e então a ideia de que robôs possam realmente ter emoções, ou que robôs possam ser criaturas vivas, acho que será difícil aceitar. Mas vamos ter que aceitar nos próximos cinquenta anos.”

Mas quanto aos robôs virem a tomar o controle, ele diz que isso provavelmente não acontecerá, por uma série de razões. Primeiro, ninguém vai construir acidentalmente um robô que queira dominar o mundo. Criar um robô que queira tomar o controle sem ter sido programado para isso seria como construir um avião 747 por acaso. E haveria tempo de sobra para impedir que isso acontecesse. Antes que alguém construa um “robô do mal”, alguém tem que construir um “robô quase mau”, e antes disso, um “robô não tão mau assim”.

Resumindo, sua filosofia é: “Os robôs estão chegando, mas não precisamos nos preocupar demais. Vai ser muito divertido.” Para ele, a revolução dos robôs é certa, e algum dia as máquinas irão superar a inteligência humana. Só resta saber quando. Mas não há nada a temer, pois os criadores somos nós. Podemos escolher criá-los para ajudar, e não atrapalhar.

HOMEM E ROBÔ SERÃO UM SÓ?

Se alguém perguntar ao dr. Brooks como iremos coexistir com robôs superinteligentes, a resposta é direta: vamos nos fundir com eles. Em vista dos avanços na robótica e nas próteses neurais, será possível incorporar a Inteligência Artificial ao nosso corpo.

Brooks observa que, em certo sentido, o processo já começou. Hoje, cerca de

20 mil pessoas têm implantes cocleares, que lhes dão o dom da audição. Os sons são captados por um minúsculo receptor que converte as ondas sonoras em sinais elétricos, que são enviados diretamente para os nervos auditivos no ouvido.

Na Universidade do Sul da Califórnia e em outras, é possível implantar uma retina artificial em pacientes cegos. Um método é colocar nos óculos uma câmara de vídeo minúscula que converte as imagens em sinais digitais. Esses sinais são enviados sem fio para um chip colocado na retina do paciente. O chip ativa os nervos da retina, que envia as mensagens do nervo óptico para o lobo occipital. Desse modo, uma pessoa totalmente cega pode enxergar uma imagem rudimentar de objetos conhecidos. Outra possibilidade é implantar na própria retina um chip sensível à luz, que envia sinais diretamente para o nervo óptico. Esse processo dispensa a câmera externa.

Isso significa também que podemos avançar ainda mais e aprimorar nossos sentidos e capacidades. O implante coclear possibilita ouvir frequências que nunca ouvimos antes. Os óculos infravermelhos já permitem ver o tipo específico de luz que emana de objetos quentes no escuro, normalmente invisível para o olho humano. A retina artificial pode aumentar nossa capacidade de ver luz ultravioleta e infravermelha. (As abelhas, por exemplo, veem luz ultravioleta porque se orientam pelo sol para voar até um canteiro de flores.)

Alguns cientistas sonham com o dia em que exoesqueletos terão superpoderes como os das histórias em quadrinhos, com superforça, supersentidos e supercapacidades. Seremos ciborgues, como o Homem de Ferro, um homem normal com superpoderes e supercapacidades. Isso significa que não precisamos nos preocupar com robôs superinteligentes tomando o poder. Vamos simplesmente nos fundir com eles.

É claro que isso está num futuro distante. Mas alguns cientistas, frustrados porque os robôs não estão saindo das fábricas para entrar em nossa vida, observam que, se a Mãe Natureza já criou a mente humana, por que não copiar? Essa estratégia consiste em desmontar o cérebro, neurônio por neurônio, e tornar a montá-lo.

Mas a engenharia reversa exige mais que um grande projeto detalhado para criar um cérebro vivo. Se um cérebro pode ser duplicado até o último neurônio, talvez possamos gravar nossa consciência num computador. Assim, poderíamos deixar para trás nosso corpo mortal. Isso vai além da questão da mente sobre o corpo. É a mente sem corpo.

4. Desde a rendição ao final da Segunda Guerra Mundial, o Japão foi forçado a se desmilitarizar, e conta hoje apenas com uma Força de Autodefesa. (N. do E.)

Prezo meu corpo, como todo mundo, mas se puder chegar aos 200 anos com um corpo de silício, eu topo.

– DANIEL HILL, COFUNDADOR DA THINKING MACHINES CORP.

11 A ENGENHARIA REVERSA DO CÉREBRO

Em janeiro de 2013, duas notícias bombásticas iriam mudar para sempre o cenário médico e científico. Da noite para o dia, a engenharia reversa do cérebro, antes considerada complexa demais para ser destrinchada, tornou-se o centro da vaidade e da rivalidade científicas entre as maiores potências econômicas do mundo.

Primeiro, em seu discurso “State of the Union”, o presidente Barack Obama deixou a comunidade científica pasma ao anunciar que uma verba federal para pesquisas, talvez da ordem de três bilhões de dólares, seria destinada ao Brain Research Through Advancing Innovative Neurotechnologies Initiative, ou projeto BRAIN. Assim como o Projeto Genoma Humano abriu as comportas da pesquisa genética, o Brain vai promover a investigação do cérebro em nível neural, mapeando os percursos elétricos. Uma vez mapeados, uma série de doenças intratáveis, como o mal de Alzheimer, de Parkinson, esquizofrenia, demência, transtorno bipolar, poderão ser entendidas e, possivelmente, curadas. Para dar a partida no projeto, foram destinados cem milhões de dólares em 2014.

Quase simultaneamente, a Comissão Europeia anunciou que o Human Brain Project teria uma verba de 1,19 bilhão de euros para criar uma simulação computadorizada do cérebro humano. Baseando-se na potência dos maiores supercomputadores do planeta, o Human Brain Project deve criar uma cópia do cérebro humano feita de transistores e aço.

Os proponentes desses dois projetos enfatizaram os enormes benefícios das empreitadas. O presidente Obama apressou-se a observar que o BRAIN não iria apenas aliviar o sofrimento de milhões de pessoas, mas também gerar novas fontes de lucro. Ele afirmou que, para cada dólar investido no Projeto Genoma Humano, foi gerada uma atividade econômica de 140 dólares. De fato, o Projeto Genoma Humano proporcionou o surgimento de várias indústrias. Para o contribuinte, o BRAIN, tal como o Projeto Genoma, só trará vantagens.

Embora Obama não tenha dado detalhes em seu discurso, os cientistas logo se manifestaram. Os neurologistas disseram que, por um lado, hoje é possível usar instrumentos delicados para monitorar a atividade elétrica de um neurônio isolado. Por outro, com aparelhos de IRM é possível monitorar o comportamento global de todo o cérebro. O que falta, dizem eles, é o campo intermediário, onde ocorre a atividade cerebral mais interessante. Nesse campo intermediário, que envolve os percursos de milhares, e até milhões, de neurônios, há enormes lacunas em nosso conhecimento sobre o comportamento e as doenças mentais.

Para lidar com esse grande problema, os cientistas lançaram um programa experimental de 15 anos de pesquisas. Nos primeiros cinco anos, os neurologistas esperam monitorar a atividade elétrica de dezenas de milhares de neurônios. As

metas em curto prazo incluem a reconstrução da atividade elétrica de partes importantes do cérebro animal, como a medula da mosca drosófila e as células ganglionares da retina de um ratinho (que tem 50 mil neurônios).

Dentro de dez anos, esse número deve chegar a centenas de milhares de neurônios. Deve incluir o mapeamento do cérebro inteiro da drosófila (135 mil neurônios) e do córtex do musaranho pigmeu, o menor mamífero conhecido, com milhões de neurônios.

Por fim, dentro de 15 anos será possível monitorar milhões de neurônios, o que corresponde ao cérebro do peixe-zebra, ou ao neocórtex inteiro de um camundongo. Isso pode abrir caminho para o mapeamento de partes do cérebro de primatas.

Enquanto isso, na Europa, o Human Brain Project adota um ponto de vista diferente para lidar com o problema. Num período de dez anos, eles usarão supercomputadores para simular o funcionamento básico de cérebros de diversos animais, começando com ratos, até chegar aos seres humanos. Em vez de lidar com neurônios isolados, o projeto usará transistores para imitar seu comportamento, de maneira que haverá módulos computacionais agindo como neocórtex, tálamo, e outras partes do cérebro.

No fim das contas, a rivalidade desses dois projetos gigantescos pode dar muitos resultados inesperados, desde descobertas para o tratamento de doenças incuráveis até o surgimento de novas indústrias. Há mais uma conquista, não verbalizada. Se algum dia for possível simular um cérebro humano, o cérebro poderá se tornar imortal? Significa que a consciência pode existir fora do corpo? Esses projetos ambiciosos trazem à luz questões muito espinhosas, teológicas e metafísicas.

CONSTRUINDO UM CÉREBRO

Assim como muitas outras crianças, eu gostava de desmontar relógios, peça por peça, tentando ver como todas se encaixavam de novo. Eu registrava mentalmente cada peça, vendo como uma engrenagem se encaixava na outra, até montar todo o relógio. Entendi que a mola mestra fazia girar a engrenagem principal, que punha em ação uma sequência de engrenagens menores, que faziam mover os ponteiros.

Hoje, em escala muito maior, cientistas da computação e neurologistas tentam desmontar um objeto infinitamente mais complexo, o mais sofisticado que temos no universo: o cérebro humano. E ainda por cima querem tornar a montá-lo, neurônio por neurônio.

Devido aos rápidos avanços nas áreas da automação, robótica, nanotecnologia

e neurociência, a engenharia reversa do cérebro humano já não é mais uma especulação inútil, tema apenas de conversas animadas depois do jantar. Nos Estados Unidos e na Europa, bilhões de dólares serão investidos em projetos que já foram considerados absurdos. Hoje, um pequeno grupo de cientistas visionários dedica sua vida profissional a um projeto que talvez eles não vivam para ver realizado. Amanhã, esse grupo poderá ser tão grande quanto um exército, generosamente financiado pelos Estados Unidos e por nações da Europa.

Se tiverem sucesso, esses cientistas poderão alterar o curso da história humana. Além de desenvolver terapias e descobrir a cura das doenças mentais, poderão revelar o segredo da consciência, e talvez gravá-la num computador.

É um projeto assustador. O cérebro humano é constituído por mais de 100 bilhões de neurônios, aproximadamente o número de estrelas que constituem a Via Láctea. Cada neurônio está conectado a talvez 10 mil outros, o que perfaz um total de 10 milhões de bilhões de conexões possíveis (e isso nem chega a computar o número de percursos existentes nesse emaranhado de neurônios). O número de “pensamentos” que um cérebro humano pode conceber é astronômico, muito além da compreensão humana.

No entanto, isso não impede que um pequeno grupo de cientistas intensamente dedicados se disponha a tentar reconstruir o cérebro a partir do zero. Um provérbio chinês diz que “Uma viagem de mil quilômetros começa com o primeiro passo”. Esse primeiro passo já foi dado pelos cientistas que decodificaram, neurônio por neurônio, o sistema nervoso de um verme nematódeo. Essa minúscula criatura, chamada *C. elegans*, tem 302 neurônios e 7 mil sinapses, todas devidamente registradas. Um esquema completo de seu sistema nervoso pode ser encontrado na internet. (Até hoje é o único organismo vivo cuja estrutura neural foi totalmente decodificada.)

A princípio, pensava-se que a engenharia reversa completa desse organismo simples abriria as portas para o cérebro humano. Ocorreu exatamente o oposto. Embora os neurônios de um nematódeo sejam finitos em número, a rede é tão complexa e sofisticada que os cientistas levaram anos para entender fatos simples sobre o comportamento desse verme, como quais percursos são responsáveis por quais comportamentos. Ao ver que até esse humilde verme nematódeo escapava ao entendimento, os cientistas foram obrigados a reconhecer a enorme complexidade do cérebro humano.

TRÊS ABORDAGENS DO CÉREBRO

Tendo em vista a complexidade do cérebro, há pelo menos três maneiras distintas

de desmontá-lo neurônio por neurônio. A primeira é simular eletronicamente o cérebro com supercomputadores, que é a abordagem adotada pelos europeus. A segunda é mapear os percursos neurais do cérebro vivo, como no projeto Brain. (Este procedimento pode ser subdividido, dependendo de como os neurônios são analisados – ou anatomicamente, neurônio por neurônio, ou por função e atividade.) Na terceira, que é uma abordagem introduzida pelo bilionário Paul Allen, da Microsoft, são decifrados os genes que controlam o desenvolvimento do cérebro.

A primeira abordagem, usando transistores e computadores para simular o órgão, está avançando com a engenharia reversa de cérebros de animais, na seguinte sequência: camundongo, rato, coelho e gato. Assim os europeus estão seguindo o acidentado caminho da evolução, começando com cérebros mais simples e prosseguindo para os mais complexos. Para um cientista da computação, a solução está na potência computacional pura – quanto mais, melhor. Isso significa usar os maiores computadores do mundo para decifrar os cérebros de ratos e homens.

O primeiro alvo é o cérebro de um camundongo, que tem um milésimo do tamanho do cérebro humano e contém cerca de 100 milhões de neurônios. O processo de pensamento no cérebro de um rato está sendo analisado pelo computador Blue Gene da IBM, no Lawrence Livermore National Laboratory, na Califórnia, onde estão alguns dos maiores computadores do mundo, usados para projetar ogivas de hidrogênio para o Pentágono. Essa quantidade colossal de transistores, chips e fios contém 147.456 processadores, com incríveis 150 mil gigabytes de memória. (Um computador doméstico comum tem um processador e alguns gigabytes de memória.)

O progresso tem sido lento, mas constante. Em vez de fazer modelos do cérebro inteiro, os cientistas tentam duplicar apenas as conexões entre o córtex e o tálamo, onde se concentra muita atividade cerebral (isso significa que faltam as conexões sensoriais com o mundo externo nessa simulação).

Em 2006, o dr. Dharmendra Modha, da IBM, usou 512 processadores para fazer uma simulação parcial do cérebro de um camundongo. Em 2007, sua equipe simulou o cérebro do rato usando 2.048 processadores. Em 2009, o cérebro de um gato, com 1,6 bilhão de neurônios e nove trilhões de conexões, foi simulado com 24.576 processadores.

Hoje, usando a potência total do computador Blue Gene, os cientistas da IBM já simularam 4,5% dos neurônios e sinapses do cérebro humano. Para dar início a uma simulação parcial do cérebro humano, são necessários 880 mil processadores, o que pode ser possível em 2020.

Tive a oportunidade de filmar o Blue Gene. Para chegar ao laboratório, passei por vários postos de segurança, pois é o principal laboratório de armas do país. E só depois desse processo, entrei numa sala imensa, climatizada, onde está o Blue

Gene.

A máquina é realmente magnífica. Consiste em vários compartimentos com enormes gabinetes pretos cheios de botões e luzes piscando, cada um com 2,5 metros de altura e cerca de 4,5 metros de comprimento. Enquanto andava entre os gabinetes que compõem o Blue Gene, eu me perguntava que operações eles estariam fazendo naquele momento. Muito provavelmente, estavam fazendo a modelagem do interior de um próton, ou calculando poços para o decaimento do plutônio, simulando a colisão de dois buracos negros ou o pensamento de um camundongo, tudo ao mesmo tempo.

Então me disseram que até aquele supercomputador estava sendo substituído pela nova geração, o Blue Gene/Q Sequoia, que levará a computação a um novo patamar. Em junho de 2012, ele estabeleceu o recorde mundial de velocidade para supercomputadores. Na velocidade de pico, ele realiza operações a 20.1 PFLOPS (ou 20.1 trilhões de operações de ponto flutuante por segundo). Ele cobre uma área de 280 metros quadrados e consome energia elétrica da ordem de 7,9 megawatts, o suficiente para iluminar uma cidade pequena.

Mas toda essa enorme potência concentrada num único computador não basta para competir com o cérebro humano?

Infelizmente, não.

As simulações por computador só tentam duplicar as interações do córtex com o tálamo. Partes enormes do cérebro ficam de fora. O dr. Modha entende a enormidade desse projeto. Suas ambiciosas pesquisas o levaram a estimar o quanto seria necessário para criar um modelo prático do cérebro humano inteiro, não apenas uma versão parcial ou simples, mas um modelo completo, com todas as partes do neocórtex e as conexões com os sentidos. Ele prevê a necessidade não de um único Blue Gene, mas de milhares deles, ocupando não só uma sala, mas um quarteirão inteiro. O altíssimo consumo de energia exige uma usina nuclear de mil megawatts para fornecer a eletricidade. E depois, para que esse computador monstruoso não derreta, é preciso desviar um rio para resfriar seus circuitos.

É incrível pensar que é preciso um computador gigantesco, quase do tamanho de uma cidade, para simular um pedaço de tecido humano que pesa menos de um quilo e meio, cabe dentro do crânio, só aumenta a temperatura do corpo em poucos graus, consome 20 watts de energia, e só precisa de alguns hambúrgueres para continuar funcionando.

CONSTRUINDO UM CÉREBRO

Talvez o cientista mais ambicioso que aderiu a esse projeto seja o dr. Henry

Markram, da École Polytechnique Fédérale de Lausanne, na Suíça. Ele é a força motriz do Human Brain Project – que recebeu financiamento de mais de um bilhão de dólares da Comissão Europeia – e passou os últimos dezessete anos tentando decodificar as conexões neurais. Ele também usa o computador Blue Gene na engenharia reversa do cérebro. No momento, a conta do Human Brain Project já está em 140 milhões de dólares, e isso representa apenas uma fração da potência computacional necessária na próxima década.

O dr. Markram acredita que este já não é mais um projeto da ciência, mas um empreendimento de engenharia, que exige muito dinheiro. Ele disse: “Para construir tudo isso – supercomputadores, softwares, pesquisa – precisamos de aproximadamente um bilhão de dólares. Não é tão caro, se considerarmos que o gasto global com as doenças do cérebro irão ultrapassar 20% do produto mundial bruto muito em breve.” Para ele, um bilhão de dólares não é nada, é uma ninharia em comparação com as centenas de bilhões gastos com mal de Alzheimer, Parkinson, e outras mazelas de aposentados.

Para o dr. Markram, o resultado será proporcional à verba. Se o projeto tiver dinheiro suficiente, o cérebro surgirá. Agora que ele conseguiu a cobiçada verba da Comissão Europeia, seu sonho poderá se tornar realidade.

Ele tem uma resposta pronta quando lhe perguntam se o contribuinte vai ser compensado por esse investimento de um bilhão, dando três razões para levar adiante essa pesquisa solitária e cara. Primeiro, “se quisermos conviver em sociedade, é essencial entender o cérebro humano, e penso que esse é um passo crucial na evolução. A segunda razão é que não podemos continuar fazendo experimentos com animais para sempre. (...) É como uma arca de Noé. É como um arquivo. E a terceira razão é que há dois bilhões de pessoas no planeta vítimas de doença mental”.

Para ele, é um absurdo sabermos tão pouco sobre doenças mentais, que causam tanto sofrimento a milhões de pessoas: “Não há uma única doença neurológica hoje em que se saiba o que está funcionando mal nos circuitos – qual o percurso, qual a sinapse, qual neurônio, qual receptor. Isso é revoltante.”

A princípio, pode parecer impossível completar esse projeto, com tantos neurônios e tantas conexões. Parece uma trabalhadeira em vão. Mas os cientistas acham que têm um ás na manga.

O genoma humano consiste em aproximadamente 23 mil genes e, no entanto, dá um jeito de criar o cérebro, que consiste em 100 bilhões de neurônios. Parece uma impossibilidade matemática criar o cérebro com esses genes, mas isso acontece cada vez que um embrião é gerado. Como tantas informações podem caber numa coisa tão pequena?

A resposta do dr. Markram é que a natureza usa atalhos. O ponto dessa abordagem é que, quando a Mãe Natureza encontra um modelo bom, certos módulos de neurônios são repetidos muitas e muitas vezes. Quando olhamos

fatias do cérebro num microscópio, a princípio só vemos um emaranhado aleatório de neurônios, mas, examinando melhor, aparecem padrões de módulos repetidos inúmeras vezes.

De fato, na arquitetura, os módulos é que possibilitam levantar grandes arranha-céus com tanta rapidez. Uma vez projetado um único módulo, é possível repeti-lo indefinidamente numa linha de montagem. Depois, basta empilhá-los, pôr um em cima do outro rapidamente, para levantar um arranha-céu. Quando toda a papelada está em ordem, em poucos meses fica pronto um prédio de apartamentos construído com módulos.

A chave do projeto Blue Brain do dr. Markram é a “coluna neocortical”, um módulo que se repete muitas vezes no cérebro. Nos humanos, cada coluna tem cerca de dois milímetros de altura com diâmetro de meio milímetro, e contém 60 mil neurônios (como ponto de comparação, cada módulo neural de um rato contém apenas 10 mil neurônios). Levou dez anos, de 1995 a 2005, para o dr. Markram mapear os neurônios numa coluna dessas e entender como ela funciona. Depois que decifrou, ele foi à IBM para criar repetições maciças dessas colunas.

Ele é um eterno otimista. Em 2009, numa conferência do TED (Tecnologia, Entretenimento e Design), ele declarou que conseguiria finalizar o projeto em dez anos. (É mais provável que seja uma versão simplificada do cérebro humano, sem as conexões dos lobos e dos sentidos.) Mas ele afirma que “se fizermos uma construção correta, ele vai falar, ter inteligência e se comportar como um humano”.

Markram é um defensor entusiasmado de seu trabalho. Tem resposta para tudo. Aos críticos que o acusam de estar adentrando um território proibido, ele responde: “Como cientistas, não precisamos ter medo da verdade. Precisamos entender nosso cérebro. É natural as pessoas pensarem que o cérebro é sagrado, que não devemos mexer nele porque pode ser que os segredos da alma estejam guardados lá. Mas, francamente, acho que, se o planeta todo entendesse como o cérebro funciona, muitos problemas seriam resolvidos. Porque então as pessoas entenderiam como os conflitos, reações e mal-entendidos são triviais, deterministas e controlados.”

Em resposta à crítica final, de que ele está “brincando de Deus”, Markram diz: “Acho que estamos longe de brincar de Deus. Deus criou todo o universo. Nós só estamos tentando construir um modelinho.”

É REALMENTE UM CÉREBRO?

Embora esses cientistas afirmem que a simulação computadorizada comece a

atingir a capacidade do cérebro humano por volta de 2020, a questão principal é: Quão realista é essa simulação? A simulação do cérebro de um gato, por exemplo, pode caçar um rato? Ou brincar com uma bolinha de lã?

A resposta é não. Essas simulações tentam se equiparar à mera potência dos neurônios ativados no cérebro do gato, mas não conseguem duplicar o modo de ligação entre as regiões do cérebro. A simulação da IBM cobre apenas o sistema tálamo-cortical (isto é, o canal que conecta o tálamo ao córtex). O sistema não tem um corpo físico, conseqüentemente, não tem as complexas interações entre o cérebro e o ambiente. Não tem lobo parietal, desse modo, não tem conexões sensoriais e motoras com o mundo externo. E mesmo dentro do sistema tálamo-cortical, as conexões básicas não reproduzem o processo de pensamento do gato. Não têm ciclos de feedback, nem circuitos de memória para caçar a presa ou encontrar um parceiro para acasalar. O cérebro computadorizado do gato é uma folha em branco, desprovida de memória e de instintos. Em outras palavras, não pode caçar um rato.

Portanto, ainda que seja possível simular o cérebro humano por volta de 2020, não poderemos ter uma conversa trivial com ele. Não tendo o lobo parietal, será como uma folha em branco, sem sensações, desprovido de qualquer conhecimento sobre si mesmo, outras pessoas e sobre o mundo à sua volta. Sem o lobo temporal, não será capaz de falar. Sem o sistema límbico, não terá emoções. De fato, terá menos aptidão cerebral que um bebê recém-nascido.

O desafio de conectar o cérebro ao mundo das sensações, emoções, linguagem e cultura está apenas começando.

ABORDAGEM POR SEGMENTAÇÃO

Outra abordagem, aprovada pelo governo de Barack Obama, é o mapeamento direto dos neurônios. Em vez de usar transistores, os próprios percursos neurais do cérebro são analisados. E há vários componentes.

Um procedimento dessa abordagem consiste em identificar fisicamente cada neurônio e cada sinapse (esse processo geralmente destrói os neurônios examinados). Chama-se abordagem anatômica. Uma alternativa é decifrar os caminhos trilhados pelos sinais elétricos entre os neurônios quando o cérebro está executando determinadas funções. Esta última, que privilegia a identificação de percursos neurais no cérebro vivo, parece ser preferida pelo governo de Obama.

A abordagem anatômica implica separar as células do cérebro de um animal, neurônio por neurônio, por um método de segmentação e análise de correlações. Assim, toda a complexidade do ambiente, o corpo, as lembranças já estão codificados no modelo. Em vez de tentar fazer um cérebro humano reunindo

uma enorme quantidade de transistores, esses cientistas querem identificar cada neurônio no cérebro. Depois, talvez cada neurônio possa ser simulado por um conjunto de transistores de modo a produzir uma réplica exata e completa do cérebro humano, com memória, personalidade e conexão com os sentidos. Após a engenharia reversa total do cérebro de alguém, será possível conversar normalmente com essa pessoa, que terá lembranças e personalidade.

Esse projeto não requer nenhuma novidade da física. O dr. Gerry Rubin, do Howard Hughes Medical Institute, usando um instrumento semelhante a uma máquina de cortar frios, já fatiou o cérebro de uma mosca-das-frutas. Não é fácil, pois o cérebro desse inseto tem apenas 300 micrômetros de espessura, um pontinho de nada em comparação com o órgão humano. O cérebro da mosca-das-frutas contém aproximadamente 150 mil neurônios. Cada fatia, da espessura de 50 bilionésimos de metro, é fotografada meticulosamente com um microscópio eletrônico, e as imagens são enviadas ao computador. Depois, um programa tenta reconstruir as conexões, neurônio por neurônio. Nesse ritmo, o dr. Rubin levará vinte anos para identificar todos os neurônios do cérebro da mosca-das-frutas.

Essa lentidão deve-se, em parte, à tecnologia fotográfica atual, pois um microscópio de varredura padrão opera a cerca de 10 milhões de pixels por segundo (isso é aproximadamente um terço da resolução de uma tela de TV por segundo). A meta é ter uma máquina de imagem capaz de processar 10 bilhões de pixels por segundo, o que seria um recorde mundial.

Armazenar os dados que saem do microscópio também é um problema. Quando o projeto desenvolver mais velocidade, Rubin espera fazer varreduras de um milhão de gigabytes de dados por dia em uma única mosca-das-frutas, armazenando-os numa quantidade de discos rígidos capaz de encher uma sala. Para completar, como cada mosca-das-frutas é ligeiramente diferente, ele terá que mapear centenas de cérebros a fim de conseguir uma aproximação exata de um deles.

Tomando por base o trabalho com cérebros dessas moscas, quanto tempo levará para fatiar o cérebro humano? “Daqui a cem anos eu gostaria de saber como funciona a consciência humana. A meta de dez ou vinte anos é entender o cérebro da mosca-da-fruta”, ele disse.

Esse método pode ser acelerado por várias técnicas avançadas. Uma possibilidade é usar um dispositivo automatizado para que o tedioso processo de fatiar e analisar cada fragmento seja realizado pela máquina. Isso reduziria muito o tempo do projeto. A automação, por exemplo, diminuiu bastante o custo do Projeto Genoma Humano (que apesar de orçado em três bilhões de dólares, foi concluído antes do prazo e abaixo do orçamento, o que é inédito em Washington). Outro método é usar uma grande variedade de tinturas para marcar diferentes neurônios e percursos, o que facilita vê-los. E uma abordagem

alternativa é criar um supermicroscópio automatizado que possa fazer varreduras extremamente detalhadas dos neurônios, um a um.

Dado que um mapeamento completo do cérebro e de todos os sentidos levará uns cem anos, esses cientistas se sentem mais ou menos como os arquitetos medievais que projetaram as catedrais europeias, sabendo que caberia a seus netos terminar a construção.

Além da construção de um mapa anatômico do cérebro, neurônio por neurônio, há um trabalho paralelo chamado “Projeto Conectoma Humano”, que usa as varreduras do cérebro para reconstruir os percursos conectando as várias regiões cerebrais.

PROJETO CONECTOMA HUMANO

Em 2010, o National Institutes of Health anunciou que destinaria trinta milhões de dólares, distribuídos ao longo de cinco anos, a um consórcio de universidades (inclusive a Universidade de Washington em Saint Louis e a Universidade de Minnesota), e 8,5 milhões distribuídos em três anos a um consórcio liderado pela Universidade de Harvard, pelo Massachusetts General Hospital e pela UCLA. Esse financiamento de curto prazo certamente não permite que os pesquisadores cheguem a um resultado completo, abrangendo o cérebro inteiro, mas o objetivo foi dar início às pesquisas.

Muito provavelmente, esse trabalho será incorporado pelo projeto BRAIN, o que vai acelerar imensamente os progressos. A meta é produzir um mapa neuronal dos percursos cerebrais, a fim de esclarecer transtornos mentais como o autismo e a esquizofrenia. Um dos líderes do Projeto Conectoma Humano, o dr. Sebastian Seung, diz: “Os pesquisadores acreditam que os neurônios em si são saudáveis, e que talvez só haja anomalias nas conexões. Mas até agora não tínhamos uma tecnologia para testar essa hipótese.” Se essas doenças são realmente causadas por más conexões, o Projeto Conectoma Humano pode nos dar pistas valiosas quanto a tratamentos.

Ao considerar o objetivo final de reunir imagens do cérebro inteiro, às vezes o dr. Seung perde a esperança de concluir o projeto. “No século XVII, o matemático e filósofo Blaise Pascal escreveu sobre seu pavor do infinito, o sentimento de insignificância ao contemplar a vastidão do espaço sideral. Como cientista, não devo falar de meus sentimentos. (...) Tenho curiosidade, fico maravilhado, mas às vezes também sinto desespero”, disse. Mas ele e outros persistem, mesmo sabendo que o projeto levará gerações para ser concluído. Não lhes faltam motivos para ter esperanças, pois algum dia microscópios automatizados tirarão fotos ininterruptamente, e máquinas com Inteligência

Artificial irão analisá-las 24 horas por dia. Hoje, porém, somente a produção de imagens de todo o cérebro humano com microscópios eletrônicos comuns consumiria um zetabyte de dados, o que é equivalente a todos os dados do mundo inteiro compilados na internet.

O dr. Seung até convida o público a participar desse grande projeto visitando o site EyeWire, onde qualquer um pode ver uma massa de percursos neurais e colori-los (respeitando seus limites). É como um livro de colorir virtual, com imagens verdadeiras de neurônios da retina gravadas por um microscópio eletrônico.

O ATLAS ALLEN DO CÉREBRO HUMANO

Por fim, há uma terceira forma de mapeamento cerebral. Em vez de analisar o cérebro por meio de simulações computadorizadas ou da identificação de todos os percursos neurais, outra abordagem recebeu a generosa verba de cem milhões de dólares do bilionário Paul Allen, da Microsoft. A meta era construir um mapa, ou atlas, do cérebro de um rato, com ênfase na identificação dos genes responsáveis pela criação do cérebro.

Espera-se que o entendimento da expressão dos genes no cérebro ajude a entender o autismo, o mal de Parkinson, de Alzheimer, e outros transtornos. Como um grande número de genes de ratos são encontrados nos humanos, é possível que os resultados nos deem uma visão mais aprofundada do cérebro humano.

Essa súbita afluência de verba levou o projeto a uma conclusão em 2006, e seus resultados estão disponíveis na internet. O Atlas Allen do Cérebro Humano foi um projeto complementar anunciado pouco depois, com o objetivo de criar um mapa anatômica e geneticamente completo do cérebro humano em 3D. Em 2011, o Instituto Allen anunciou que tinha mapeado a bioquímica de dois cérebros humanos, encontrando mil locais anatômicos com cem milhões de pontos de dados, detalhando como os genes se expressam na bioquímica subjacente do cérebro. O estudo confirmou que 82% de nossos genes estão expressos no cérebro.

“Até então, simplesmente não existia um mapa definitivo do cérebro humano com esse nível de detalhe”, disse o dr. Allen Jones, do Instituto Allen. “O Atlas Allen do Cérebro Humano oferece imagens jamais vistas do mais complexo e mais importante dos nossos órgãos.”

Os cientistas que dedicaram a vida à engenharia reversa do cérebro sabem que têm décadas de muito trabalho pela frente. Mas também conhecem as implicações práticas de seu trabalho. Sabem que mesmo os resultados parciais ajudarão a decodificar os mistérios das doenças mentais que sempre afligiram a humanidade.

Os céticos, porém, podem alegar que após a conclusão dessa tarefa árdua teremos uma montanha de dados e nenhum entendimento de como se encaixam. Por exemplo: imagine um homem de Neandertal encontrando um diagrama completo do computador Blue Gene da IBM. Todos os detalhes, até o último transistor, estão no imenso diagrama, com milhares de metros quadrados de papel. Esse homem de Neandertal pode ter uma vaga noção de que aquilo é o segredo de uma máquina superpoderosa, mas aquela massa de dados técnicos não significa nada para ele.

Da mesma forma, há o receio de que, após bilhões de dólares gastos para decifrar a localização de cada neurônio no cérebro, não sejamos capazes de entender o que aquilo tudo significa. Pode levar muitas décadas de trabalho árduo para vermos como o todo funciona.

Por exemplo: o Projeto Genoma Humano foi um sucesso tremendo no sequenciamento de todos os genes que compõem o genoma humano, mas uma grande decepção para quem esperava a cura imediata de doenças genéticas. O Projeto Genoma Humano é um dicionário gigantesco, com 23 mil verbetes, mas nenhuma definição. São páginas e páginas inócuas, embora a “ortografia” de cada gene seja perfeita. O projeto de fato rompeu uma barreira, mas ao mesmo tempo é apenas o primeiro passo numa longa jornada para entender o que os genes fazem e como eles interagem.

Da mesma forma, ter um mapa completo de todas as conexões neurais no cérebro não garante entendermos o que os neurônios fazem e como reagem. A engenharia reversa é a parte fácil. Depois, vem a parte difícil – compreender todos os dados.

O FUTURO

Suponhamos que, finalmente, o momento tenha chegado. Com muita festa, os cientistas anunciam que realizaram a engenharia reversa do cérebro humano inteiro.

E aí?

Uma aplicação imediata é encontrar as origens de certas doenças mentais. Acredita-se que muitas doenças mentais não são causadas pela destruição dos neurônios, mas por uma simples conexão defeituosa. Pense em doenças

genéticas causadas por uma simples mutação, como a doença de Huntington, de Tay-Sachs, ou a fibrose cística. Dentre três bilhões de pares de bases, um único errinho (ou repetição) pode causar movimentos involuntários dos membros e convulsões, como na doença de Huntington. Ainda que o genoma esteja 99,9999999% correto, uma falha mínima pode invalidar a sequência inteira. É por isso que a terapia dos genes considera essas simples mutações como o possíveis doenças genéticas que podem ser curadas.

Do mesmo modo, após a engenharia reversa do cérebro, será possível fazer simulações, rompendo intencionalmente algumas conexões para verificar se provocam certas doenças. Uns poucos neurônios podem ser responsáveis por transtornos graves de cognição. Uma função da engenharia reversa pode localizar esse pequeno grupo de neurônios mal ativados.

Um exemplo é o delírio de Capgras, no qual o paciente pode ver a própria mãe, mas achar que se trata de uma impostora. Segundo o dr. V. S. Ramachandran, esse transtorno raro pode ser causado por má conexão entre duas partes do cérebro. O giro fusiforme, no lobo temporal, é responsável pelo reconhecimento do rosto da mãe, e a amígdala é responsável pela reação emocional ao ver a mãe. Quando esses dois centros estão desconectados, o indivíduo reconhece perfeitamente o rosto da mãe, mas, como não há reação emocional, acha que ela é uma impostora.

Outra aplicação da engenharia reversa do cérebro é localizar exatamente o grupo de neurônios que está falhando. A estimulação cerebral profunda, como vimos, envolve o uso de sondas minúsculas para amortecer uma parte minúscula do cérebro, como a área 25 de Broadmann, no caso de certas formas graves de depressão. Com o mapa da engenharia reversa é possível saber exatamente onde os neurônios estão falhando, mesmo que isso envolva apenas um pequeno grupo deles.

A engenharia reversa do cérebro pode ser de grande ajuda também para a Inteligência Artificial. A visão e o reconhecimento de rostos são tarefas muito fáceis para o cérebro, mas ainda escapam aos mais avançados computadores. Por exemplo: os computadores podem reconhecer com exatidão de 95%, ou mais, um rosto que faça parte de um pequeno banco de dados, mas, se o mesmo rosto for mostrado em ângulos diferentes, ou se não estiver no banco de dados, o computador provavelmente não vai reconhecer. Nós podemos reconhecer rostos familiares, por diversos ângulos, em um décimo de segundo. É tão fácil para o cérebro que nem percebemos o que estamos fazendo. A engenharia reversa pode revelar o mistério de como isso é feito.

Mais complicadas são as doenças que envolvem falhas cerebrais múltiplas, com a esquizofrenia. Esse distúrbio abrange diversos genes, além de interações com o ambiente, o que, por sua vez, provoca uma atividade anormal em várias áreas do cérebro. Mas, mesmo nesses casos, a engenharia reversa pode apontar

exatamente onde certos sintomas (por exemplo, alucinações) são formados, o que pode talvez indicar um caminho para a cura.

A engenharia reversa do cérebro poderia também solucionar problemas básicos, mas não resolvidos, como o modo de armazenamento da memória de longo prazo. Sabe-se que certas partes do cérebro, como o hipocampo e a amígdala, armazenam memória, mas não se sabe como a memória é distribuída pelos vários córtices e depois reorganizada para criar uma lembrança.

Quando a engenharia reversa do cérebro for plenamente funcional, virá o momento de ligar todos os circuitos para ver se respondem como um humano (isto é, se passam no teste de Turing). Como a memória de longo prazo já está codificada nos neurônios que passaram pela engenharia reversa, deve ser fácil verificar rapidamente se o cérebro pode responder de modo indistinguível ao de um humano.

Por fim, há um impacto da engenharia reversa no cérebro que raramente é mencionado, mas está na mente de muitos: a imortalidade. Se a consciência pode ser transferida para um computador, isso significa que não morreremos?

Especulação não é perda de tempo. Ela abre os caminhos entre os galhos secos do matagal da dedução.

– ELIZABETH PETERS

Somos uma civilização científica. (...) Isso significa uma civilização em que o conhecimento e sua integridade são cruciais. Ciência não é nada mais que uma palavra que vem do latim e quer dizer conhecimento. (...)

Conhecimento é o nosso destino.

– JACOB BRONOWSKI

12 O FUTURO: A MENTE ALÉM DA MATÉRIA

A consciência pode existir por si mesma, livre das restrições do corpo físico? Podemos deixar nosso corpo e, como os espíritos, vagar por esse parque de diversões chamado universo?

Isso foi explorado em *Jornada nas estrelas*, quando o capitão Kirk, da nave Enterprise, encontra uma raça super-humana, quase um milhão de anos mais avançada que a Federação dos Planetas. São tão avançados que há muito tempo abandonaram seus frágeis corpos e agora habitam globos pulsantes de pura energia. Há milênios que não têm sensações inebriantes, como respirar ar fresco, tocar a mão de alguém, ou receber carinho físico. O líder, Sargon, dá boas-vindas à Enterprise em seu planeta. O capitão Kirk aceita o convite, ciente de que aquela civilização poderia desintegrar sua nave instantaneamente, se quisesse.

Mas a tripulação não sabe que aqueles superseres têm uma fraqueza fatal. Apesar de toda a tecnologia que possuem, estão há centenas de milhares de anos privados do corpo físico. Sendo assim, sentem falta das sensações físicas e querem voltar a ser humanos.

Um desses superseres é mau e está decidido a se apossar do corpo físico de um membro da tripulação. Ele quer viver como humano, ainda que isso destrua a mente do dono do corpo. A entidade má se apossa do corpo de Spock e desencadeia uma batalha com a tripulação no deque da Enterprise.

Os cientistas se perguntam: Existe uma lei da física contra a existência da mente fora do corpo? Mais especificamente, se a mente humana é um aparelho que cria continuamente modelos de mundo e os simula no futuro, é possível produzir uma máquina capaz de simular todo esse processo?

Já mencionamos a possibilidade de termos nosso corpo colocado num casulo, como no filme *Substitutos*, enquanto controlamos mentalmente um robô. O problema aqui é que nosso corpo natural vai continuar definhando, enquanto nossos substitutos, robôs, continuarão existindo. Cientistas sérios estão considerando a possibilidade de transferirmos a mente para um robô, e nos tornarmos realmente imortais. E quem não gostaria de ter a chance de viver para sempre? Como disse Woody Allen: “Não quero viver para sempre por meio da minha obra. Quero viver para sempre sem morrer.”

Milhões de pessoas afirmam que é possível a mente sair do corpo. De fato, muitos alegam ter conseguido.

EXPERIÊNCIAS FORA DO CORPO

A ideia de uma mente fora do corpo é talvez uma de nossas crenças mais antigas,

integrada em mitos, folclore, sonhos, e talvez até nos genes. Toda sociedade tem histórias de espíritos e demônios que entram e saem do corpo humano quando bem entendem.

Infelizmente, muitos inocentes foram perseguidos para que os demônios que supostamente habitavam seu corpo fossem exorcizados. Eles provavelmente sofriam de alguma doença mental, como a esquizofrenia, em que a pessoa é atormentada por vozes geradas no próprio cérebro. Os historiadores acreditam que uma das bruxas enforcadas em Salém, em 1692, provavelmente tinha uma doença genética rara, chamada mal de Huntington, que causa agitação incontrolável dos membros.

Hoje, algumas pessoas afirmam ter entrado num estado de transe em que a consciência deixa o corpo e vagueia pelo espaço, capaz até de ver o próprio corpo de longe. Numa enquete com 13 mil europeus, 5,8% afirmaram ter tido uma experiência fora do corpo. Entrevistas com pessoas nos Estados Unidos apresentam números similares. O ganhador do Prêmio Nobel de Física Richard Feynman, sempre curioso a respeito de novos fenômenos, certa vez se enfiou num tanque de isolamento, com total privação sensorial, para tentar sair do corpo físico. E conseguiu. Depois relatou que saiu do corpo, vagou pelo espaço, e viu o próprio corpo inerte. No entanto, mais tarde ele concluiu que o episódio provavelmente foi obra de sua imaginação devido à privação sensorial.

Neurologistas que estudaram esse fenômeno têm uma explicação mais prosaica. Na Suíça, o dr. Olaf Blanke e seus colegas podem ter identificado no cérebro o local exato que gera as experiências fora do corpo. Uma paciente, de 43 anos, sofria de convulsões que se originavam no lobo temporal direito. Colocaram uma touca com 100 eletrodos sobre a sua cabeça para localizar a região responsável pelas convulsões. Quando os eletrodos estimularam a área entre os lobos parietal e temporal, ela imediatamente teve a sensação de sair do corpo. “Eu me vejo de cima, deitada na cama, mas só vejo as pernas e a parte de baixo do tronco!”, exclamou. Ela sentiu estar flutuando dois metros acima do corpo.

Quando os eletrodos foram desligados, a sensação desapareceu instantaneamente. O dr. Blanke concluiu que podia ligar e desligar a sensação de estar fora do corpo, como num interruptor de luz, estimulando repetidamente essa área cerebral. Como vimos no capítulo 9, lesões epiléticas no lobo temporal podem induzir o sentimento de que há espíritos maus por trás de qualquer infelicidade, de modo que o conceito do espírito saindo do corpo pode ser parte da nossa constituição neural.

Isso explicaria também a presença de seres sobrenaturais. Quando o dr. Blanke analisou uma paciente de 22 anos que sofria de convulsões intratáveis, concluiu que, ao estimular a área temporoparietal, podia induzir a sensação de que havia uma presença sinistra atrás dela. A paciente conseguia descrever em detalhes

essa pessoa, que chegou até a agarrar seus braços. A posição do vulto mudava a cada aparição, mas sempre surgia atrás dela.

Acredito que a consciência humana é o processo de formar continuamente um modelo do mundo, a fim de simular o futuro e atingir um objetivo. O cérebro recebe sensações principalmente dos olhos e do ouvido interno para criar um modelo de onde nos situamos no espaço. No entanto, quando os sinais dos olhos e dos ouvidos entram em contradição, o cérebro fica confuso quanto à localização. Geralmente, temos náuseas e vomitamos. Por exemplo: muitos sentem enjoo a bordo de um navio porque seus olhos, fixos nas paredes da cabine, lhes dizem que estão parados, enquanto a orelha interna lhes diz que estão balançando. A incongruência dos sinais provoca náusea. O remédio é olhar para o horizonte, pois assim a imagem visual é compatível com os sinais da orelha interna. (Essa mesma sensação de náusea pode ser induzida até mesmo se estivermos parados. Se você ficar olhando para uma lata de lixo girando, pintada com listras coloridas verticais, as listras parecem se mover horizontalmente, dando a sensação de que você está se movendo, mas sua orelha interna diz que você está parado. A incongruência faz com que a pessoa vomite, mesmo se estiver sentada numa cadeira.)

As mensagens dos olhos e do ouvido interno podem também ser distorcidas eletricamente, na fronteira entre os lobos temporal e parietal, e essa é a origem da experiência fora do corpo. Quando essa área tão sensível é tocada, o cérebro fica confuso a respeito de sua localização no espaço. (Vale notar que uma perda temporária de sangue, de oxigênio, ou excesso de dióxido de carbono na corrente sanguínea também podem causar uma interrupção na região temporal-parietal e induzir experiências fora do corpo. Talvez isso possa explicar a prevalência de tais sensações em casos de acidentes, emergências, infartos etc.)

EXPERIÊNCIAS DE QUASE MORTE

Talvez a categoria mais intensa de experiências fora do corpo esteja nas histórias de indivíduos que foram declarados mortos e, misteriosamente, recobram a consciência. De 6% a 12% de sobreviventes de ataques cardíacos relatam ter tido a experiência de quase morte. É como se tivessem enganado a morte. Em entrevistas, eles contam uma história dramática da mesma experiência: saíram do corpo e foram flutuando em direção a uma luz brilhante no fim de um longo túnel.

A mídia se valeu desses relatos para gerar inúmeros best-sellers e documentários de TV dedicados a essas histórias espetaculares. Muitas teorias estranhas foram propostas para explicar a experiência de quase morte. Numa

enquete com duas mil pessoas, 42% disseram acreditar que tal experiência é uma prova do contato com o mundo espiritual que existe além da morte. (Alguns acreditam que o corpo libera endorfinas – narcóticos naturais – antes de morrer. Isso pode explicar a euforia que as pessoas sentem, mas não o túnel e a luz brilhante.) Carl Sagan chegou a especular que a experiência de quase morte seria um reviver do trauma do nascimento. O fato desses indivíduos relatarem sempre a mesma experiência não confirma necessariamente que tocaram na vida após a morte. Na verdade, parece indicar a ocorrência de algum evento neurológico profundo.

Neurologistas que examinaram seriamente esse fenômeno suspeitam que o motivo pode ser a diminuição do fluxo de sangue no cérebro, que acompanha frequentemente esses casos, e que também ocorre em desmaios. O dr. Thomas Lempert, neurologista da Castle Park Clinic, em Berlim, conduziu uma série de experimentos, provocando desmaios em 42 indivíduos em condições controladas de laboratório. Destes, 60% tiveram alucinações visuais (por exemplo, luzes brilhantes e manchas coloridas), 47% disseram ter entrado em outro mundo, 20% afirmaram ter encontrado um ser sobrenatural, 17% viram uma luz brilhante, e 8% viram um túnel. Portanto, o desmaio pode induzir as sensações que as pessoas têm em experiências de quase morte. Mas, como isso acontece?

A explicação para um desmaio simular a experiência de quase morte pode ser formulada a partir de uma análise de experiências com pilotos das Forças Armadas. A Força Aérea dos Estados Unidos, por exemplo, chamou o neurofisiologista dr. Edward Lambert para analisar pilotos militares que desmaiaram quando expostos a valores altos de força g (por exemplo, ao executar uma curva muito acentuada num jato, ou subindo depois de mergulhar muito baixo).^[5] O dr. Lambert colocou os pilotos numa ultracentrífuga na Mayo Clinic, em Rochester, Minnesota, que girava até atingir uma grande aceleração. Passados 15 segundos, o sangue foi drenado do cérebro e os pilotos ficaram inconscientes.

O dr. Lambert concluiu que depois de cinco segundos de aceleração o fluxo de sangue nos olhos dos pilotos diminuía, e a visão periférica se estreitava, criando a imagem de um longo túnel. Isso explicaria o túnel visto por pessoas que tiveram experiência de quase morte. Se a visão periférica é apagada, só se vê um longo túnel à frente. Como o dr. Lambert era capaz de ajustar a velocidade da centrífuga, percebeu que poderia manter os pilotos naquele estado indefinidamente, demonstrando que a visão do túnel é causada pela falta de fluxo sanguíneo na periferia dos olhos.

Alguns cientistas que investigaram as experiências de quase morte e de sair do corpo estão convencidos de que elas são subproduto do próprio cérebro quando colocado em condições estressantes e suas conexões ficam confusas. Contudo, outros cientistas acreditam que daqui a décadas, quando nossa tecnologia estiver avançada o suficiente, a consciência poderá realmente sair do corpo. Alguns métodos, controversos, foram sugeridos.

Um método foi lançado pelo inventor e futurista dr. Ray Kurzweil, que acredita que a consciência poderá um dia ser gravada num supercomputador. Certa vez, fizemos uma apresentação numa conferência e ele me falou que sua fascinação por computadores e Inteligência Artificial começou aos 5 anos de idade, quando seus pais compravam todo tipo de brinquedos e aparelhinhos automáticos. Ele adorava montar e desmontar aqueles brinquedos e, já criança, sabia que estava destinado a ser um inventor. Fez doutorado no MIT, sob orientação do dr. Marvin Minsky, um dos pioneiros da IA. Em seguida, dedicou-se a aplicar a tecnologia de reconhecimento de padrões em instrumentos musicais e máquinas de transposição de texto para som. Sua pesquisa sobre IA nessas áreas lhe rendeu uma série de empresas. (Aos 20 anos ele já tinha vendido sua primeira empresa.) Seu leitor óptico, que reconhece e converte texto em som, foi anunciado como um auxílio para cegos, chegando a ser mencionado por Walter Cronkite no noticiário da noite.

Ele me disse que para ter sucesso como inventor é preciso estar sempre além da curva, prever as mudanças, e não apenas reagir a elas. De fato, o dr. Kurzweil adora fazer previsões, e muitas delas espelham o crescimento exponencial da tecnologia digital. Ele fez as seguintes previsões:

- Em 2019, um microcomputador de mil dólares terá a potência computacional de um cérebro humano – 20 milhões de bilhões de cálculos por segundo. (Esse número foi obtido tomando-se os 100 bilhões de neurônios do cérebro multiplicados por mil conexões por neurônio e 200 cálculos por segundo, por conexão.)
- Em 2029, um microcomputador de mil dólares será mil vezes mais potente que o cérebro humano; a engenharia reversa do próprio cérebro humano terá êxito.
- Em 2055, mil dólares em potência computacional serão iguais à potência de processamento conjunto de todos os humanos do planeta. (Modestamente, ele acrescentou: “Posso estar errando por um ano ou dois.”)

Para Kurzweil, o ano de 2045 será, particularmente, muito importante, pois ele

acredita que a “singularidade” estará estabelecida. As máquinas terão superado os humanos em inteligência, criando novas gerações de robôs mais inteligentes do que elas. Ainda segundo o dr. Kurzweil, dado que esse processo pode continuar indefinidamente, haverá um aumento constante na potência das máquinas. Nesse cenário, será melhor nos fundirmos com nossas criações, ou pelo menos não atrapalhar. (Embora essas datas estejam num futuro distante, ele me disse que quer viver para ver o dia em que os humanos finalmente serão imortais, ou seja, ele quer viver o bastante para viver para sempre.)

Como sabemos, a partir da Lei de Moore, em determinado ponto a potência computacional não vai mais conseguir avançar criando transistores cada vez menores. Na opinião de Kurzweil, a única maneira de continuar expandindo a potência computacional será aumentar o tamanho das máquinas, o que levará os robôs a devorar todos os minerais da Terra para alimentar sua potência cada vez maior. Quando o planeta tiver se tornado um computador gigante, os robôs poderão ser obrigados a explorar o espaço sideral em busca de novas fontes de potência computacional. Podem chegar a consumir a energia de estrelas inteiras.

Uma vez perguntei a ele se esse crescimento cósmico dos computadores poderia alterar o próprio cosmos. Ele respondeu que sim. Contou que às vezes fica contemplando o céu à noite, imaginando se em algum planeta distante seres inteligentes já atingiram a singularidade. Se sim, talvez tenham deixado nas estrelas alguma marca visível a olho nu.

Uma limitação, disse ele, é a velocidade da luz. Se as máquinas não conseguirem romper essa barreira, o aumento exponencial da potência vai esbarrar num teto. Quando isso acontecer, talvez as próprias leis da física sejam alteradas.

Quem faz previsões com tamanha exatidão e extensão naturalmente atrai críticas como um para-raios, mas ele não se abala. As pessoas podem resmungar, contestando essa ou aquela previsão porque Kurzweil errou alguns prazos, mas ele está mais interessado no desenvolvimento de ideias que preveem o crescimento exponencial da tecnologia. Na verdade, muitas pessoas que entrevistei na área da IA concordam que alguma forma de singularidade irá acontecer, mas discordam quanto às datas e desdobramentos. Por exemplo: Bill Gates, cofundador da Microsoft, acredita que nenhum contemporâneo viverá para ver o dia em que os computadores serão inteligentes o suficiente para se passar por humanos. Kevin Kelly, editor da revista *Wired*, disse: “As pessoas que preveem um futuro muito utópico, sempre preveem que acontecerá antes de morrerem.”

Na verdade, um dos muitos objetivos de Kurzweil é trazer seu pai de volta à vida. Ou melhor, ele deseja criar uma simulação realista. Há várias possibilidades, mas todas são altamente especulativas.

Kurzweil propõe que talvez o DNA de seu pai possa ser extraído (do túmulo, de

parentes, ou de materiais orgânicos que ele tenha deixado). Contendo mais ou menos 23 mil genes, o DNA seria um diagrama completo para recriar o corpo e produzir um clone do indivíduo.

Certamente, é uma possibilidade. Uma vez perguntei ao dr. Robert Lanza, da Advanced Cell Technology, como ele conseguiu trazer “de volta à vida” uma criatura morta há muito tempo, o que fez história. Ele me disse que o zoológico de San Diego lhe pedira para criar um clone de um banteng, um animal parecido com um boi, que tinha morrido vinte e cinco anos antes. Foi difícil extrair uma célula válida para criar um clone, mas ele conseguiu. Depois enviou a célula por correio expresso para uma fazenda e a célula foi implantada numa vaca, que deu à luz o animal. Quanto a nenhum primata ter sido ainda clonado, muito menos um humano, Lanza acha que é um problema técnico, e que é apenas uma questão de tempo para que consigam clonar um ser humano.

Mas essa é a parte fácil. O clone será geneticamente equivalente ao original, mas sem a memória. Lembranças artificiais podem ser gravadas no cérebro do clone com os métodos avançados descritos no capítulo 5, como inserção de sondas no hipocampo, ou a criação de um hipocampo artificial, mas o pai de Kurzweil morreu há muito tempo, e é impossível fazer uma gravação. O melhor que se pode fazer é juntar fragmentos de todos os dados históricos da pessoa, entrevistando quem lembre de fatos relevantes, acessando suas transações de cartão de crédito etc., e inserindo as informações no programa.

Um modo mais prático de inserir a personalidade e memória de alguém é criar um grande arquivo de dados contendo todas as informações conhecidas sobre os hábitos e a vida da pessoa. Por exemplo: hoje é possível guardar todos os e-mails, transações com cartões de crédito, anotações, agendamentos, diários eletrônicos e história de vida num único arquivo, que pode criar uma imagem exata de quem você é. O arquivo representará sua “assinatura digital”, contendo todo o conhecimento sobre você. E pode ser muito minucioso e íntimo, mostrando quais são seus vinhos preferidos, como você passou as férias, que sabonete você usa, seu cantor predileto, e muito mais.

Utilizando um questionário, é possível ter uma noção aproximada da personalidade do pai de Kurzweil. Amigos, parentes e colegas poderiam responder ao extenso questionário sobre a personalidade dele, se era tímido, curioso, honesto, trabalhador etc., atribuindo graus a cada traço de personalidade (por exemplo, “10” para muito honesto). Isso criaria uma lista com centenas de números, cada um atribuindo um valor a um traço específico. Quando todos esses números são compilados, um programa de computador usa as informações para dar um resultado aproximado do comportamento dele em situações hipotéticas. Digamos que você esteja dando uma palestra e seja confrontado por algum mal-educado. O programa de computador pode pesquisar os números e prever um dentre vários comportamentos possíveis (por exemplo, você pode

ignorar o desbocado, xingar de volta, partir para a briga). Em outras palavras, a personalidade básica do pai de Kurzweil seria reduzida a uma longa lista de números, de 1 a 10, usados pelo computador para prever como ele reagiria em diversas situações.

O resultado seria um vasto programa de computador que responderia a situações mais ou menos como a pessoa em questão, usando as mesmas expressões verbais, com as mesmas peculiaridades e idiossincrasias, tudo ajustado levando em conta as lembranças sobre a pessoa.

Outra possibilidade é abandonar o processo de clonagem e criar um robô parecido com o original. Basta inserir esse programa num sócia mecânico que fala como o indivíduo original, com o mesmo sotaque, os mesmos maneirismos, e faz os mesmos movimentos. E acrescentar expressões recorrentes dele (por exemplo, “entendeu?”) também seria fácil.

Naturalmente, com a tecnologia que temos hoje seria óbvio perceber que se trata de um robô. No entanto, nas próximas décadas será possível se aproximar cada vez mais do original, até que fique bom o suficiente para enganar os desavisados.

Isso levanta uma questão filosófica. Essa “pessoa” será mesmo igual à original? A original continua morta, de modo que o clone ou robô é, no sentido exato, um impostor. Um gravador, por exemplo, pode reproduzir uma conversa com total fidelidade, mas certamente o gravador não é a pessoa original. Um clone ou robô que se comporte exatamente como o original pode ser um substituto válido?

IMORTALIDADE

Esses métodos têm sido criticados porque o processo não grava realisticamente personalidade e memória verdadeiras. Um modo mais confiável de colocar a mente na máquina é por meio do Projeto Conectoma Humano, discutido no capítulo anterior, que pretende duplicar, neurônio por neurônio, todos os percursos celulares do cérebro. Todas as lembranças e esquisitices da personalidade estão inseridas no conectoma.

O dr. Sebastian Seung, do Projeto Conectoma, observa que há quem pague 100 mil dólares ou mais para manter o cérebro congelado em nitrogênio líquido. Certos animais, como peixes e sapos, podem passar o inverno congelados num bloco de gelo e continuar perfeitamente saudáveis na primavera, quando descongelam. Isso porque eles usam a glucose para alterar o ponto de congelamento da água no sangue. Assim, o sangue permanece líquido, mesmo com o bicho preso num bloco de gelo. No corpo humano, porém, essa alta

concentração de glicose provavelmente seria fatal; portanto, congelar um cérebro humano em nitrogênio líquido é um procedimento duvidoso, porque a expansão dos cristais de gelo vai romper a parede celular de dentro para fora (e também, quando as células cerebrais morrem, são invadidas por íons de cálcio, causando a expansão delas até o rompimento). De qualquer forma, é muito provável que as células cerebrais não sobrevivam ao processo de congelamento.

Em vez de congelar o corpo e romper as células, um processo mais confiável para obter a imortalidade talvez seja completar o conectoma da pessoa. Nesse caso, o médico terá todas as conexões neurais do paciente num disco rígido. Basicamente, a alma fica num disco, reduzida a informações. No futuro, alguém será capaz de ressuscitar o conectoma e, em princípio, usar um clone ou um monte de transistores para trazer a pessoa de volta à vida.

O Projeto Conectoma, como dissemos, ainda está longe de poder gravar as conexões neurais humanas. Mas, como diz o dr. Seung, “Devemos ridicularizar os pesquisadores da imortalidade, chamando-os de tolos? Ou algum dia eles irão rir no nosso túmulo?”

DOENÇA MENTAL E IMORTALIDADE

A imortalidade pode ter suas desvantagens. Os cérebros eletrônicos construídos até agora contêm apenas as conexões entre o córtex e o tálamo. O cérebro gerado por engenharia reversa, não tendo corpo, pode sofrer de isolamento sensorial e até manifestar sinais de doença mental, como os prisioneiros confinados na solitária. Talvez o preço de se criar um cérebro imortal, com engenharia reversa, seja a loucura.

Cobaias humanas colocadas em câmaras de isolamento, privadas de qualquer contato com o mundo externo, acabam alucinando. Em 2008, a BBC exibiu um programa científico chamado *Total Isolation*, em que foram monitorados seis voluntários confinados num bunker nuclear, sozinhos e em completa escuridão. Passados dois dias, três deles começaram a ver e ouvir coisas – cobras, carros, zebras e ostras. Quando foram libertados, os médicos constataram que todos eles tinham sofrido deterioração mental. A memória de um deles teve uma queda de 36%. Pode-se imaginar que, após semanas ou meses, a maioria teria enlouquecido.

Para manter a sanidade de um cérebro construído por engenharia reversa, será essencial conectá-lo a sensores que recebam sinais do ambiente exterior, que ele tenha a possibilidade de ver e ter sensações do mundo externo. Mas então surge outro problema: o cérebro pode se sentir como uma aberração grotesca, uma cobaia esdrúxula à mercê de um experimento científico. Como o cérebro

mantém a memória e a personalidade do humano original, anseia por contato humano. Mas, espiando lá dentro da memória de um supercomputador, com seu emaranhado de eletrodos e fios pendurados, o cérebro de engenharia reversa seria repulsivo para qualquer humano. Seria impossível se afeiçoar a ele. Os amigos lhe dariam as costas.

O PRINCÍPIO DO HOMEM DAS CAVERNAS

Nesse ponto, surge o que chamo de Princípio do Homem das Cavernas. Por que tantas previsões razoáveis falham? E por que alguém *não* quereria viver para sempre dentro de um computador?

O Princípio do Homem das Cavernas é o seguinte: podendo escolher entre contato tecnológico e contato pessoal, sempre optamos pelo pessoal. Por exemplo: se pudermos escolher entre ganhar ingressos para um show ao vivo do nosso cantor predileto ou um CD do mesmo cantor, o que escolhemos? Ou, se pudermos escolher entre ganhar uma passagem de avião para visitar o Taj Mahal ou uma belíssima foto do palácio, o que vamos preferir? É mais que provável escolhermos os ingressos para o show e a passagem de avião.

Isso acontece porque herdamos a consciência de nossos ancestrais macacos. É provável que algo de nossa personalidade básica não tenha mudado muito nos últimos 100 mil anos, quando os primeiros humanos surgiram na África. Grande parte da nossa consciência é dedicada a manter uma boa aparência e a tentar impressionar nossos amigos e os membros do sexo oposto.

O mais provável é que, dada nossa consciência básica semelhante à dos macacos, só iremos nos fundir com computadores se eles fizerem melhorias, mas não substituírem totalmente nosso corpo atual.

O Princípio do Homem das Cavernas deve explicar por que algumas previsões razoáveis sobre o futuro não se materializaram, como o “escritório sem papel”. Os computadores deveriam extinguir o papel nos escritórios, mas, ironicamente, até criaram mais papéis. Isso porque descendemos de caçadores, que precisam “provar que mataram, mostrando a caça” (isto é, confiamos na evidência concreta, não em elétrons efêmeros dançando num computador e que somem quando o desligamos). Da mesma forma, a “cidade sem gente”, onde os habitantes usariam a realidade virtual para participar de reuniões em vez de ter que sair de casa, não se materializou. Ir e vir nas cidades está cada vez mais complicado. Por quê? Porque somos animais sociais, que gostam de conviver com outros. A videoconferência, apesar de muito útil, não capta todo o espectro de informações sutis oferecidas pela linguagem corporal. Se um chefe quer investigar um problema da equipe, por exemplo, ele quer ver os subordinados

tremendo e suando para responder ao interrogatório. Só se pode fazer isso pessoalmente.

O HOMEM DAS CAVERNAS E A NEUROCIÊNCIA

Quando criança, li a trilogia *Fundação*, de Isaac Asimov, que me influenciou profundamente. Primeiro, me obrigou a formular uma pergunta simples: Como será a tecnologia daqui a 50 mil anos, quando houver um império galáctico? E ao longo do livro, tive que me perguntar: Por que os humanos têm a mesma aparência e agem da mesma maneira que hoje? Eu achava que dali a milhares de anos, certamente os humanos teriam corpos de ciborgue com capacidade super-humana. Teriam largado seus corpos insignificantes milênios atrás.

Encontrei duas respostas. Primeiro, Asimov queria conquistar um público jovem, então criou personagens com quem esse público se identificasse, inclusive com todas as falhas deles. Segundo, talvez as pessoas no futuro tivessem a opção de ter um corpo com superpoderes, mas preferissem parecer normais na maior parte do tempo. Isso porque a mente delas não mudou desde o surgimento dos seres humanos, e a aceitação dos outros e do sexo oposto ainda determinava sua aparência e o que desejavam da vida.

Agora, vamos aplicar o Princípio do Homem das Cavernas à neurociência do futuro. Significa, no mínimo, que qualquer modificação da forma humana básica teria que ser quase invisível exteriormente. Não queremos parecer saídos de cenas de filmes de ficção científica, com eletrodos e fios pendurados na cabeça. Os implantes cerebrais para inserir memória ou desenvolver a inteligência só serão adotados se, com a nanotecnologia, surgirem sondas e sensores microscópicos invisíveis a olho nu. No futuro, será possível produzir nanofibras, feitas talvez de nanotubos de carbono com uma molécula de espessura, capazes de atingir os neurônios com precisão cirúrgica e aprimorar as capacidades mentais sem interferir na aparência.

Enquanto isso, se precisarmos estar conectados a um supercomputador para fazer o *upload* de informações, não queremos ficar presos a um cabo enfiado na coluna, como em *Matrix*. A conexão deve ser sem fio, para que possamos acessar grandes quantidades de potência computacional apenas localizando mentalmente o servidor mais próximo.

Hoje temos implantes cocleares e retinas artificiais que conferem os dons da audição e da visão a pacientes, mas, no futuro, nossos sentidos serão aprimorados pela nanotecnologia, preservando nossa forma humana básica. Por exemplo: podemos ter a opção de melhorar os músculos via modificação genética ou exoesqueletos. Haverá também lojas de corpo humano, onde poderemos

encomendar peças de reposição quando as antigas estiverem gastas, mas esses e outros aprimoramentos do corpo devem evitar a perda da forma humana.

Outra maneira de aplicar essa tecnologia conforme o Princípio do Homem das Cavernas é usá-la como opção, e não como um modo permanente de vida. Podemos preferir a opção de nos conectar à tecnologia e desconectar pouco depois. Cientistas podem querer potencializar a inteligência para resolver um problema particularmente difícil, e depois tirar o capacete ou os implantes e ir cuidar de outras coisas. Assim não precisamos aparecer como soldadinhos espaciais na frente de nossos amigos. O objetivo é que nada disso seja feito por obrigação. Precisamos ter a opção de usufruir dos benefícios da nova tecnologia sem o vexame de parecer meio bobos.

Nos séculos que estão por vir, é provável que nosso corpo seja muito semelhante ao que temos hoje, exceto que será aperfeiçoado e terá maiores poderes. O fato de nossa consciência ser dominada por desejos e anseios antigos é um vestígio do nosso passado simiesco.

Mas e a imortalidade? Como vimos, um cérebro gerado por engenharia reversa, com todas as idiossincrasias do indivíduo original, acabará ficando louco se colocado dentro de um computador. Além disso, conectar esse cérebro a sensores externos para ter as sensações do ambiente é criar um monstro grotesco. Uma solução parcial é conectá-lo a um exoesqueleto. Se o exoesqueleto age como substituto, o cérebro da engenharia reversa pode ter sensações como toque e visão, sem a aparência grotesca. Depois, quando for criado o exoesqueleto sem fio, poderá agir como humano, mas será controlado pelo cérebro da engenharia reversa “vivo” dentro do computador.

Esse substituto terá o melhor dos dois mundos. Sendo um exoesqueleto, ele será perfeito. Possuirá superpoderes. Conectado sem fio ao cérebro criado pela engenharia reversa dentro de um grande computador, ele será imortal. E como poderá sentir o ambiente e ser bonito como um humano verdadeiro, não terá grandes problemas na interação com outros humanos, já que muitos também já terão optado pelo mesmo procedimento. Assim, o verdadeiro conectoma fica num supercomputador fixo, enquanto sua consciência se manifesta no belo corpo ambulante de um substituto.

Tudo isso exige um nível de tecnologia muito além da que temos hoje. Entretanto, em vista da rapidez dos progressos científicos, pode se tornar realidade antes do fim do século.

TRANSFERÊNCIA GRADUAL

Hoje, o processo de engenharia reversa requer a transferência das informações

que estão dentro do cérebro, neurônio por neurônio. O cérebro tem que ser cortado em fatias muito finas, pois a varredura por IRM ainda não está aperfeiçoada a ponto de identificar com total precisão a arquitetura neural do cérebro vivo. Portanto, até que isso aconteça, a desvantagem óbvia do procedimento é que precisamos morrer para nos submeter à engenharia reversa. Como o cérebro degenera rapidamente após a morte, a preservação tem que ser providenciada imediatamente, o que é muito difícil de se conseguir.

Mas pode haver um meio de atingir a imortalidade sem ter que morrer antes. Essa ideia foi lançada pelo dr. Hans Moravec, ex-diretor do Laboratório de Inteligência Artificial da Universidade Carnegie Mellon. Quando o entrevistei, ele disse que antevê um futuro distante, quando será possível fazer a engenharia reversa do cérebro com o propósito específico de transferir a mente para um corpo robótico imortal, mesmo com o indivíduo consciente. Se for possível fazer a engenharia reversa de cada neurônio no cérebro, por que não criar uma cópia, feita com transistores, duplicando exatamente o processo de pensamento? Assim não será preciso morrer para viver eternamente. Poderemos ficar conscientes durante todo o processo.

Ele me disse que esse processo tem que ser feito passo a passo. Primeiro, você se deita numa maca ao lado de um robô sem cérebro. Em seguida, um robô-cirurgião extrai alguns neurônios do seu cérebro e os duplica em alguns transistores situados no robô. Seu cérebro é conectado por fios aos transistores na cabeça vazia do robô. Os neurônios são jogados fora e substituídos pelo circuito de transistores. Como o cérebro ainda está ligado pelos fios aos transistores, continua funcionando normalmente, e você permanece consciente durante o processo. Então o supercirurgião continua removendo os neurônios, sempre duplicando-os nos transistores do robô. A meio caminho da operação, metade do seu cérebro está vazia. A outra metade está conectada por fios aos transistores da cabeça do robô. Quando todos os neurônios tiverem sido removidos, o indivíduo terá um cérebro-robô que será uma duplicata exata do seu cérebro original, neurônio por neurônio.

Quando o processo chega ao fim, você se levanta da maca e se vê num corpo perfeito. Você tem o corpo e o rosto mais belos do que jamais sonhou, com capacidades e poderes super-humanos. E ainda tem a regalia da imortalidade. Ao olhar para o corpo original, vê apenas uma carcaça envelhecida e descerebrada.

É claro que essa tecnologia está muito além do nosso tempo. Não podemos nem aplicar engenharia reversa no cérebro humano, quanto mais produzir uma cópia exata feita de transistores. Uma das principais críticas a essa abordagem é que um cérebro transistorizado não cabe dentro de um crânio. Dado o tamanho dos componentes eletrônicos, um cérebro transistorizado deve ser do tamanho de um supercomputador enorme. Nesse sentido, essa proposta remete à anterior,

em que o cérebro de engenharia reversa é armazenado num imenso supercomputador que, por sua vez, controla um substituto. Mas a grande vantagem dessa abordagem é que não será preciso morrer, e continuaremos conscientes durante todo o processo.

Essas possibilidades chegam a dar tontura. Todas parecem ter coerência com as leis da física, mas as barreiras tecnológicas são enormes. Todas as propostas de gravação da consciência num computador requerem uma tecnologia que está num futuro longínquo.

Contudo, uma última proposta para chegar à imortalidade não exige nenhuma engenharia reversa do cérebro. Exige apenas um microscópico “nanorrobô” capaz de manipular átomos um a um. Então, por que não viver para sempre em nosso corpo natural, só com um “ajuste” periódico para mantê-lo imortal?

O QUE É ENVELHECER?

Essa nova abordagem incorpora a pesquisa mais recente sobre o envelhecimento. Tradicionalmente, não há consenso entre biólogos sobre a fonte desse processo. Mas na última década, uma teoria foi ganhando maior aceitação, unindo diversas vias de pesquisa sobre a velhice. Basicamente, o envelhecimento é um acúmulo de erros em nível genético e celular. À medida que as células envelhecem, erros vão se acumulando no DNA e armazenam-se nas células, tornando-as vagarosas. As células começam a funcionar mal, a pele vai ficando flácida, os ossos se enfraquecem, os cabelos começam a cair, o sistema imunológico se deteriora. Um dia, morremos.

As células também contam com mecanismos de correção de erros, mas com o passar do tempo o dispositivo começa a falhar, e o envelhecimento se acelera. A meta, portanto, é fortalecer os mecanismos naturais de reparo celular, o que pode ser feito pela terapia dos genes e a criação de novas enzimas. Existe, porém, outro meio: o “nanorrobô” montador.

Uma das peças dessa tecnologia futurista é algo chamado “nanorrobô”, uma máquina atômica minúscula que faz o patrulhamento da corrente sanguínea, eliminando células cancerosas, reparando os danos do processo de envelhecimento e nos mantendo jovens e saudáveis para sempre. A natureza já criou alguns nanorrobôs na forma de células imunes que patrulham o corpo através do sangue. Essas células atacam vírus e corpos estranhos, mas não o processo de envelhecimento.

Se os nanorrobôs puderem reverter os estragos do envelhecimento em nível molecular e celular, a imortalidade estará ao nosso alcance. Nesse cenário, os nanorrobôs são como células imunes, policiais minúsculos patrulhando a corrente

sanguínea. Eles atacam as células cancerosas, neutralizam vírus, eliminam detritos e mutações. Assim a possibilidade de ser imortal estará ao alcance usando nosso próprio corpo, sem precisar de robô nem de clone.

NANORROBÔS – REALIDADE OU FANTASIA?

Tenho a filosofia de que se algo é compatível com as leis da física, os únicos impedimentos para torná-lo realidade são da ordem da engenharia e da economia. As barreiras econômicas e de engenharia são grandes, mas mesmo impraticáveis no momento, são transponíveis.

À primeira vista, o nanorrobô é simples: uma máquina com braços e cortadores que agarram e partem as moléculas em pontos específicos, e depois as juntam de novo. Cortando e colando vários átomos, o nanorrobô pode criar quase todas as moléculas conhecidas, como um mágico tirando algo da cartola. E também pode se autorreproduzir: basta construir um único exemplar. Este digere as matérias-primas e cria milhões de outros nanorrobôs. Isso pode desencadear uma segunda Revolução Industrial, pois o custo do material de construção vai despencar. Talvez um dia todas as casas tenham seu próprio montador molecular, e será possível ter tudo o que quiser. Basta pedir a ele.

Mas a questão principal é: os nanorrobôs são coerentes com as leis da física? Em 2001, dois visionários travaram uma briga pública sobre essa questão. O que estava em jogo era uma concepção de todo o futuro da tecnologia. De um lado, o falecido Richard Smalley, ganhador do Prêmio Nobel de Química, era cético quanto aos nanorrobôs. Do outro lado estava Eric Drexler, um dos pais da nanotecnologia. A batalha, titânica e interminável, ocupou páginas de várias revistas científicas de 2001 a 2003.

Smalley dizia que, em escala atômica, surgem novas forças quânticas que tornam impossível a construção de nanorrobôs. O erro de Drexler e de outros, dizia ele, é que o nanorrobô, com braços e cortadores, não pode funcionar em escala atômica. Existem novas forças (por exemplo, a força de Casimir) que fazem os átomos se repelirem ou se atraírem. Ele se referiu a isso como um problema de “dedos grossos e grudentos” porque os dedos do nanorrobô não são pinças e chaves delicadas e precisas. Forças quânticas atrapalham; é como tentar soldar metais usando luvas muito grossas. E cada vez que tentamos juntar as duas peças de metal, elas se repelem ou grudam na gente; não é possível pegá-las da forma adequada.

Drexler contra-atacou, dizendo que nanorrobôs não são ficção científica, que realmente existiam. Propôs que ele pensasse nos ribossomos do nosso corpo. Tais estruturas são essenciais para criar e fundir as moléculas de DNA; cortam e

juntam moléculas em pontos específicos, o que possibilita a formação de novas cadeias de DNA.

Smalley não se convenceu. Afirmou que os ribossomas não são máquinas que conseguem cortar e colar o que nós quisermos. Eles agem especificamente em moléculas de DNA. Além disso, os ribossomas são feitos de substâncias químicas orgânicas, que precisam de enzimas para acelerar a reação, o que ocorre apenas em ambiente aquoso. E concluiu dizendo que os transistores são feitos de silício, não de água, portanto as enzimas jamais funcionariam. Drexler rebateu afirmando que catalisadores podem funcionar sem água. A briga continuou acalorada, em vários rounds. No fim, como competidores empatados, os dois estavam exaustos. Drexler admitiu que a analogia com cortadores e maçaricos era simplista demais, e forças quânticas podiam interferir. E Smalley reconheceu não ter conseguido nocautear o adversário, pois a natureza tinha pelo menos uma forma de solucionar o problema dos “dedos grossos e grudentos”, com ribossomas, e talvez haja outros meios mais sutis, ainda desconhecidos.

Evitando entrar nos detalhes desse debate, Ray Kurzweil acredita que os nanorrobôs, com ou sem dedos grossos e grudentos, um dia irão moldar não só moléculas, mas a própria sociedade. Ele resumiu sua visão ao dizer: “Não estou planejando morrer. (...) Vejo isso, em última análise, como um despertar de todo o universo. Penso que hoje o universo inteiro é feito basicamente de energia e matéria burras, mas que irá despertar. E se despertar transformado nessa energia e matéria de sublime inteligência, espero fazer parte dele.”

Por mais fantásticas que sejam essas especulações, não passam de um prefácio para a próxima rodada de especulações. Talvez um dia a mente não só estará livre do corpo material, mas também será capaz de explorar o universo, em forma de energia pura. A ideia de que a consciência um dia será livre para vagar pelas estrelas é o sonho máximo. E, por incrível que pareça, está em perfeita conformidade com as leis da física.

5. Força g é uma medida da aceleração sentida por uma pessoa como se fosse seu peso. Na superfície da Terra sentimos uma força de $1g$, que é o nosso próprio peso. (N. do R.T.)

13 A MENTE COMO ENERGIA PURA

A ideia de que um dia a consciência possa se espalhar pelo universo foi seriamente considerada por vários físicos. Sir Martin Rees, o Astrônomo Real da Grã-Bretanha, escreveu: “Buracos de minhoca, dimensões extras e computadores quânticos abrem cenários especulativos que podem transformar o universo num ‘cosmos vivo!’”

Mas, a mente algum dia se libertará do corpo material para explorar o universo todo? Esse foi o tema abordado por Isaac Asimov no clássico *A última pergunta* (ele recorda afetuosamente que, de todos os contos de ficção científica que escreveu, esse é seu favorito). No conto, bilhões de anos no futuro, os humanos já tinham deixado o corpo físico em casulos num planeta obscuro, libertando a mente para controlar a energia pura em toda a galáxia. Seus substitutos não eram de aço e silício, mas seres de pura energia que percorriam facilmente todos os confins do espaço, passando por explosões de estrelas, colisões de galáxias, e outras maravilhas do universo. Contudo, por mais poderosa que a humanidade tivesse se tornado, estava impotente diante da morte definitiva do universo, o Grande Congelamento. Desesperados, os humanos construíram um supercomputador que respondesse à pergunta final: A morte do universo pode ser revertida? O computador era tão grande e complexo que precisou ser colocado no hiperespaço, mas só conseguia dizer que os dados eram insuficientes para dar a resposta.

Muitas eras depois, quando as estrelas começaram a escurecer e toda a vida no universo estava à beira da morte, o supercomputador finalmente descobriu um meio de reverter o fim do universo. Ele coletou todas as estrelas mortas do universo, fundiu-as numa gigantesca bola cósmica e a incendiou. Enquanto a bola explodia, o supercomputador anunciou: “Faça-se a luz!”

E fez-se a luz.

Assim a humanidade, liberta do corpo físico, foi capaz de brincar de Deus e criar um novo universo.

Em princípio, o fantástico conto de Asimov, com seres de energia pura vagando pelo universo, parece impossível. Estamos habituados a pensar em seres de carne e osso, sujeitos às leis da física e da biologia, vivendo e respirando na Terra, presos à gravidade do planeta. O conceito de entidades de energia conscientes vasculhando as galáxias, livres das limitações do corpo material, é muito estranho.

No entanto, o sonho de explorar o universo como seres de energia pura cabe muito bem nas leis da física. Pense na forma mais conhecida de energia pura, o raio laser, que é capaz de conter grandes quantidades de informação. Hoje, trilhões de sinais, na forma de telefonemas, pacotes de dados, vídeos e mensagens por e-mail são transmitidos continuamente por cabos de fibra óptica

conduzindo raios laser. Um dia, talvez no próximo século, seremos capazes de transmitir a consciência do cérebro através do sistema solar colocando nossos conectomas inteiros em potentes raios laser. Um século depois disso, poderemos enviar nosso conectoma às estrelas, viajando num raio laser. Isso é possível porque o comprimento de onda do raio laser é microscópico, isto é, medido em milionésimos de metro, o que significa que podemos comprimir grandes quantidades de informação nesse padrão ondulatório. Pense no Código Morse. Os traços e pontos do código podem ser facilmente superpostos ao padrão de onda do raio laser. E mais informações ainda podem ser transferidas num feixe de raios X, que tem um comprimento de onda menor que um átomo.

Um meio de explorar a galáxia, livre das restrições e do estorvo da matéria, é colocar nosso conectoma em raios laser direcionados para a Lua, os planetas, e até mesmo para as estrelas. Tendo o programa prioritário para identificar os percursos cerebrais, o conectoma completo do cérebro humano estará disponível até o fim deste século e, no próximo século, poderemos ter uma forma de conectoma contida num raio laser.

O raio laser deve conter todas as informações necessárias para reconstituir um ser consciente. Embora possa levar anos, ou mesmo séculos, até que o raio laser chegue ao seu destino, do ponto de vista de quem estiver se deslocando pelo laser, a viagem é instantânea. A consciência ficaria basicamente congelada no raio laser ao atravessar o espaço, de modo que a viagem até o outro lado da galáxia parece se dar num piscar de olhos.

Desse modo, evitamos todos os inconvenientes de viagens interplanetárias e interestelares. Primeiro, não é preciso construir foguetes de propulsão colossais. Basta ligar o botão do laser. Segundo, não há poderosas forças g esmagando o corpo na aceleração rumo ao espaço. O indivíduo é lançado imediatamente na velocidade da luz, pois seu estado é imaterial. Terceiro, ele não fica exposto aos perigos do espaço sideral, como impactos de meteoros e raios cósmicos letais, pois os asteroides e a radiação o atravessam, sem causar danos. Quarto, não é preciso congelar o corpo, nem passar anos de tédio em câmara lenta num foguete convencional. É só zunir pelo espaço na maior velocidade do universo, congelado no tempo.

Ao chegar ao destino, há uma estação de recepção para transferir os dados do raio laser para um computador central, que traz a consciência de volta à vida. O código impresso no raio laser assume o controle do computador e o reprograma. O conectoma leva o computador central a iniciar o processo de simulação do futuro a fim de alcançar seus objetivos (isto é, se tornar consciente).

O ser consciente no computador central envia sinais sem fio para um corpo robótico substituto que já está à espera no local de destino. Assim, “despertamos” subitamente num planeta ou estrela distante, dentro do corpo robótico do substituto, como se a viagem tivesse se passado num piscar de olhos. Todos os

cálculos complexos ocorrem no grande computador central, que direciona os movimentos do substituto para continuar sua missão numa estrela distante. E nem ficamos sabendo dos percalços da viagem espacial; é como se nada tivesse acontecido.

Agora imagine uma grande rede de estações de recepção espalhadas pelo sistema solar e pela galáxia. Desse ponto de vista, pular de estrela em estrela seria fácil, viajando à velocidade da luz, em jornadas instantâneas. Em cada estação há um substituto robótico à espera para entrarmos em seu corpo, como um quarto vazio de hotel reservado para nós. Chegamos ao destino, descansados e equipados com um corpo super-humano.

O tipo de substituto robótico à nossa espera ao fim da viagem depende da missão. Se for explorar um mundo novo, vai precisar de um corpo para trabalhar em condições adversas. É preciso ser ajustado para um campo gravitacional diferente, atmosferas tóxicas, temperaturas congelantes ou escaldantes, ciclos de dia e noite diferentes, e uma chuva constante de radiação letal. Para sobreviver a essas condições extremas, o corpo substituto precisa ter superforça e supersentidos.

Se o corpo substituto for apenas para o lazer, deve ser projetado para essas atividades, maximizando o prazer de voar pelo espaço em esquis, pranchas de surfe, parapentes, asas delta, avião, ou de jogar bola no espaço, com bastão, taco ou raquete.

E se tiver a missão de estudar e interagir com povos de outros planetas, o substituto precisa ter as características corporais da população nativa (como no filme *Avatar*).

Para criar uma rede de estações de laser, antes de tudo é preciso viajar para os planetas e estrelas à moda antiga, em naves espaciais convencionais. E só depois será possível construir as primeiras estações de laser.

Talvez o meio mais rápido, barato e eficaz de implantar uma rede interestelar seja enviar sondas robóticas autorreplicantes para explorar a galáxia. Como essas sondas são capazes de fazer cópias de si mesmas, basta começar com uma delas, e depois de muitas gerações haverá bilhões de sondas cruzando o universo em todas as direções, cada uma criando uma estação de laser onde quer que aterrissem. Discutiremos isso no próximo capítulo.

Quando a rede estiver totalmente instalada, podemos pensar num fluxo contínuo de seres conscientes viajando pela galáxia, com multidões chegando e partindo a todo momento. Uma estação da rede será algo parecido com a Grand Central Station de Nova York.

Por mais futurista que pareça, a física básica desse conceito já está bem estabelecida. É possível colocar grandes quantidades de dados em raios laser, enviar essas informações para milhares de quilômetros de distância e decodificá-las no ponto de destino. Os principais problemas que essa ideia enfrenta não estão

na física, mas na engenharia. Por isso, talvez só no próximo século seja possível enviar nosso conectoma inteiro em raios laser com potência suficiente para alcançar outros planetas. E pode levar outro século até o laser conduzir nossa mente às estrelas.

Para verificar a viabilidade disso, é necessário fazer alguns cálculos simples e rápidos. O primeiro problema é que os fótons dentro de um raio laser, extremamente finos, embora pareçam estar numa formação perfeitamente paralela, divergem ligeiramente no espaço. Quando criança, eu acendia uma lanterna na direção da Lua, perguntando-me se a luz poderia chegar lá. A resposta é sim. A atmosfera absorve mais de 90% do fecho original, deixando um resto que chega à Lua. Mas o verdadeiro problema é que a imagem que a lanterna finalmente projeta na Lua se estende por quilômetros. Isso ocorre devido ao princípio de incerteza; até os raios laser divergem com o tempo. Como não podemos saber a localização precisa do raio laser, então, pelas leis da física quântica, ele deve se espalhar lentamente no correr do tempo.

Mas levar nosso conectoma até a Lua não é grande vantagem. É muito mais simples e fácil ficar na Terra, controlando o substituto lunar diretamente por rádio. O atraso no envio de ordens para o substituto é de aproximadamente um segundo. A maior vantagem do laser está no controle de substitutos nos planetas, pois uma mensagem por rádio pode levar horas para chegar lá. O processo de enviar uma série de ordens por rádio ao substituto, esperar a resposta, e enviar outra ordem é lento demais, e pode levar dias.

Para enviar raios laser aos planetas, o primeiro passo é estabelecer uma bateria de lasers na Lua, bem acima da atmosfera, onde não tem ar para absorver o sinal. Disparado da Lua, o raio laser pode chegar a um planeta em questão de minutos, ou algumas horas. Quando o laser tiver levado o conectoma ao planeta, será possível controlar diretamente o substituto, sem qualquer demora.

Portanto, a instalação de uma rede de estações de laser ligando o sistema solar pode ser concretizada no próximo século. Os problemas são muito maiores para enviar raios laser às estrelas. É preciso ter estações de retransmissão em asteroides e estações espaciais em todo o caminho para amplificar o sinal, reduzir erros de transmissão, e enviar a mensagem à estação seguinte. Potencialmente, isso pode ser feito usando os cometas que ficam entre o Sol e as estrelas mais próximas. Por exemplo: aproximadamente a um ano-luz do Sol (um quarto da distância até a estrela mais próxima), estão os cometas da chamada nuvem de Oort. É uma concha esférica de bilhões de cometas, muitos dos quais ficam parados no espaço vazio. Provavelmente, há uma nuvem de cometas semelhante à de Oort em volta das estrelas do sistema Alfa Centauro, que é o vizinho estelar mais perto de nós. Supondo que essa outra nuvem de Oort esteja também a um ano-luz dessas estrelas, então metade da distância entre

nosso sistema solar e o próximo sistema solar contém cometas estacionários, onde podemos construir estações de retransmissão de laser.

Outro problema é a quantidade de dados a ser enviada por raios laser. O total de informações contidas num conectoma, segundo o dr. Sebastian Seung, é por volta de um zetabyte (um *zeta* é 1 seguido por 21 zeros). É mais ou menos o equivalente ao total de informações contidas hoje na internet. Agora, imagine disparar no espaço vários raios laser carregando essa montanha de informações. As fibras óticas podem transmitir terabytes de dados por segundo (um *tera* é 1 seguido por 12 zeros). No próximo século, os avanços no armazenamento de informação, compressão de dados e feixes de raios laser podem aumentar essa eficiência um milhão de vezes. Isso significa que em poucas horas todas as informações contidas num cérebro humano podem cruzar o espaço, transportadas por raios laser.

Então o problema não se resume à quantidade de dados transmitidos por laser. Em princípio, os raios laser podem transportar uma quantidade ilimitada de dados. O gargalo fica nas estações de recepção, dos dois lados, que precisam ter mecanismos para manipular essa quantidade de dados a uma velocidade espantosa. Transistores de silício talvez não tenham velocidade suficiente para lidar com esse volume de dados. É preciso ter computadores quânticos, que não usam transistores de silício, mas átomos, um a um. Hoje, os computadores quânticos estão num nível muito primitivo, mas daqui a um século devem ter potência suficiente para lidar com zetabytes de informação.

SERES FLUTUANTES DE ENERGIA

Outra vantagem dos computadores quânticos no processamento dessa imensa quantidade de dados é a possibilidade de criar seres de energia capazes de flutuar pelo ar, como costumamos ver em cenas de ficção científica e fantasia. Esses seres representam a consciência na forma pura. Em princípio, eles parecem violar as leis da física, pois a velocidade da luz é sempre a mesma.

Mas na última década, quem virou manchete foram os físicos da Universidade de Harvard ao anunciar que conseguiram parar completamente um raio de luz. Aparentemente, eles realizaram o impossível, desacelerando a luz até que ficasse lenta ao ponto de conseguirem engarrafá-la. Prender um raio de luz numa garrafa não é assim tão fantástico, se olharmos atentamente para um copo de água. Quando a luz penetra na água, desacelera, curvando-se quando entra num ângulo inclinado. Da mesma forma, a luz se curva quando penetra no vidro, o que possibilita a existência de microscópios e telescópios. A razão disso é dada pela teoria quântica.

Pense no antigo Pony Express, o serviço de correio do século XIX no Oeste dos Estados Unidos. O cavalo percorria com velocidade a distância entre os postos de muda. O gargalo era devido à demora nos postos, onde se dava a troca de cavalo, cavaleiro e correspondência. Isso diminuía consideravelmente a velocidade média do serviço. Da mesma forma, no vácuo entre os átomos, a luz ainda viaja com velocidade c , a velocidade da luz, que é aproximadamente 299.792 quilômetros por segundo. Entretanto, quando atinge os átomos, a luz sofre um atraso, pois é rapidamente absorvida e reemitida, retomando a trajetória uma fração de segundo depois. Essa breve demora é responsável pelos raios de luz, na média, aparentemente diminuir a velocidade no vidro ou na água.

Os cientistas de Harvard exploraram esse fenômeno, resfriando um recipiente com gás à temperatura de quase zero absoluto. Nessas baixíssimas temperaturas, os átomos absorvem um raio de luz por períodos muito mais longos antes de reemitir-lo. Assim, aumentando o fator de retardo, eles conseguiram desacelerar o raio até que ficasse em repouso. O raio de luz continuava viajando à velocidade da luz entre os átomos de gás, mas passava um tempo cada vez maior sendo absorvido por eles.

Isso abre a possibilidade para que um ser consciente, em vez de controlar um substituto, prefira permanecer na forma de energia pura e viajar pelo espaço, como um fantasma.

No futuro, quando raios laser com nossos conectomas forem enviados às estrelas, um feixe pode ser transferido para uma nuvem de moléculas de gás, e então ser colocado numa garrafa. A “garrafa de luz” é muito semelhante a um computador quântico. Ambos têm um conjunto de átomos vibrando em uníssono, em que os átomos estão em fase uns com os outros. E ambos podem fazer cálculos complexos, muito além da capacidade de um computador normal. Pois bem, se os problemas dos computadores quânticos forem resolvidos, também poderemos ter a capacidade de lidar com “garrafas de luz”.

MAIS RÁPIDO QUE A LUZ?

Como vimos, todos esses problemas são da área da engenharia. Não há lei da física que impeça de viajarmos num fecho de energia no próximo século, ou depois. Talvez esse seja o meio mais conveniente de visitar planetas e estrelas. Em vez de pilotar um raio de luz, como sonham os poetas, podemos ser o raio de luz.

Para captar bem a ideia do conto de ficção científica de Asimov, devemos perguntar se a viagem intergaláctica mais rápida que a luz é realmente possível. No conto, seres de imenso poder se movimentam livremente entre galáxias

separadas por milhões de anos-luz.

Será possível? Para responder a essa pergunta, temos que empurrar as fronteiras da física quântica moderna. Basicamente, os chamados “buracos de minhoca” podem ser um atalho na vastidão do espaço e do tempo. E seres de energia pura, e não de matéria, terão uma vantagem decisiva para atravessá-los.

Em certo sentido, Einstein é como um policial de plantão, afirmando que não se pode viajar mais rápido que a luz, que é a maior velocidade possível no universo. Viajar pela nossa galáxia, a Via Láctea, por exemplo, levaria cem mil anos, mesmo a bordo de um raio laser. Embora só se passasse um instante para o passageiro, o tempo no planeta nativo teria progredido cem mil anos. E uma viagem entre galáxias levaria de milhões a bilhões de anos-luz.

Mas o próprio Einstein deixou uma brecha em sua obra. Na teoria da relatividade geral, de 1915, ele mostrou que a gravidade surgia da deformação do espaço-tempo. A gravidade não é o “puxão” de uma misteriosa força invisível, como pensava Newton, mas um “empurrão” causado pelo próprio espaço, que se curva em torno de um objeto. Além dessa ideia ser uma explicação brilhante para a luz estelar se curvar ao passar perto das estrelas, e para a expansão do universo, abre a possibilidade de o tecido do espaço-tempo ir se esgarçando até rasgar.

Em 1935, Einstein e seu aluno Nathan Rosen introduziram a possibilidade de dois buracos negros se juntarem pelas costas, grudados como irmãos siameses. Desse modo, ao cair num buraco negro, seria factível, em princípio, passar para o outro. (Imagine juntar dois funis pelo furo estreito de cada um. A água que escorre por um funil sai pelo outro.) O “buraco de minhoca”, também chamado de Ponte de Einstein-Rosen, introduziu a possibilidade de portais, ou passagens, entre universos. O próprio Einstein reconheceu a inviabilidade de passarmos por um buraco negro, pois seríamos esmigalhados nesse processo. Mas vários avanços subsequentes levantaram a possibilidade de uma viagem mais rápida que a luz, pegando o atalho de um buraco de minhoca.

Em 1963, o matemático Roy Kerr descobriu que um buraco negro girando não colapsa em um simples ponto, como se pensava antes, mas se torna um anel rotatório, tão rápido que as forças centrífugas impedem seu colapso. Se você cair nesse anel, pode passar para outro universo. As forças gravitacionais são enormes, mas não infinitas. Como em *Alice através do espelho*, você passa a mão pelo espelho e entra num universo paralelo. A borda desse “espelho” é o anel que forma o buraco negro. Desde a descoberta de Kerr, muitas outras soluções das equações de Einstein mostram que, em princípio, é possível atravessar universos sem ser esmigalhado instantaneamente. Como todos os buracos negros já vistos no espaço giram em alta velocidade (alguns foram cronometrados a 1,6 milhão de quilômetros por hora), os portais cósmicos talvez sejam comuns.

Em 1988, o físico dr. Kip Thorne, da Cal Tech, e seus colegas mostraram que, com “energia negativa” suficiente, pode ser possível estabilizar um buraco negro, de modo que um buraco de minhoca se torne “atravessável” (isto é, pode-se passar por ele em qualquer direção sem ser esmagado). A energia negativa é a substância talvez mais exótica do universo, mas realmente existe e pode ser criada (em quantidades microscópicas) em laboratório.

Então surge um novo paradigma. Primeiro, uma civilização avançada deve concentrar energia positiva suficiente num único ponto, comparável à de um buraco negro, para abrir uma passagem pelo espaço ligando dois pontos distantes. Segundo, deve acumular energia negativa suficiente para manter o portal aberto, para que se mantenha estável e não se feche no momento em que você entrar.

Agora, vamos pôr essa ideia nas devidas proporções. No fim deste século, deve ser possível mapear todo o conectoma humano. Uma rede interplanetária para lasers poderá ser implantada no começo do próximo século, de modo que a consciência possa ser transmitida e atravessar o sistema solar. Não será preciso haver nenhuma nova lei da física. Uma rede de laser capaz de passar pelas estrelas talvez precise esperar até outro século. Mas uma civilização que use livremente buracos de minhoca estará milhares de anos à nossa frente em tecnologia, ampliando as fronteiras da física conhecida.

Tudo isso tem implicações diretas para viabilizar ou não a passagem da consciência por entre universos. Se a matéria chega perto de um buraco negro, a gravidade fica tão intensa que o corpo vira “espaguete”. A gravidade que incide sobre a perna é maior que a gravidade que puxa a cabeça, e o corpo é esticado por forças de maré. De fato, na proximidade de um buraco negro, os átomos do corpo são esticados até que os elétrons sejam arrancados do núcleo, e os átomos se desintegram.

Para imaginar o poder das forças de maré, pense nas marés da Terra e nos anéis de Saturno. A gravidade da Lua e a do Sol dão um puxão sobre a Terra que faz as águas dos mares subirem alguns metros na maré cheia. E se uma lua chegar perto demais de um planeta enorme como Saturno, as forças de maré vão esticar essa lua até dilacerá-la. A distância em que uma lua é dilacerada pelas forças de maré é chamada de limite de Roche. Os anéis de Saturno ficam exatamente no limite de Roche, portanto devem ter sido criados por uma lua que se aproximou demais do planeta-mãe.

Se entrarmos num buraco negro em rotação, mesmo usando a energia negativa para estabilizá-lo, os campos gravitacionais podem ser tão fortes que viramos espaguete.

É aqui que o raio laser tem uma vantagem importante sobre a matéria ao passar por um buraco de minhoca. Como a luz laser é imaterial, não pode ser esticada por forças de maré ao passar perto de um buraco negro. Em vez disso, a luz sofre um “desvio para o azul” (isto é, ganha energia e sua frequência

aumenta). Mesmo com o raio laser distorcido, a informação contida nele permanece intacta. Por exemplo: uma mensagem em código Morse transmitida por raio laser fica comprimida, mas a informação permanece igual. A informação digital é intocada por forças de maré. Assim, a força gravitacional, fatal para seres feitos de matéria, é inofensiva para seres viajando em raios de luz.

Dessa maneira, a consciência conduzida por raio laser, por ser imaterial, tem uma vantagem decisiva sobre a matéria ao passar por um buraco de minhoca.

Os feixes de laser têm ainda outra vantagem sobre a matéria. Alguns físicos calcularam que pode ser mais fácil criar um buraco de minhoca microscópico, talvez do tamanho de um átomo. A matéria não poderia passar por um buraco tão minúsculo, mas o laser de raio-X, com comprimento de onda menor que um átomo, seria capaz de passar sem dificuldade.

Embora o genial conto de Asimov seja claramente uma obra da imaginação, ironicamente, talvez já exista na galáxia uma rede de estações de laser. No entanto, somos tão primitivos que nem sabemos disso. Para uma civilização milhares de anos à nossa frente, a tecnologia de digitalização de conectomas e envio às estrelas deve ser brincadeira de criança. Nesse caso, é concebível que seres inteligentes já estejam lançando suas consciências através de uma vasta rede de raios laser pela galáxia. Nem os nossos telescópios e satélites mais avançados conseguem detectar uma rede intergaláctica dessas.

Carl Sagan lamentava a possibilidade de estarmos cercados de civilizações alienígenas e não termos tecnologia para perceber.

A próxima questão é: O que se esconde na mente alienígena?

Se encontrarmos uma civilização tão avançada, que tipo de consciência ela terá? Um dia, o destino da raça humana pode depender dessa resposta.

Às vezes penso que o sinal mais seguro de que existe vida inteligente em outros lugares do universo é que nunca tentaram fazer contato conosco.

– BILL WATTERSON

Ou existe vida inteligente no espaço, ou não existe. As duas alternativas são assustadoras.

– ARTHUR C. CLARKE

14 A MENTE ALIENÍGENA

Em *A guerra dos mundos*, de H. G. Wells, marcianos atacam a Terra porque seu planeta está morrendo. Armados de raios mortíferos e gigantescas máquinas que andam, eles incineram rapidamente várias cidades e estão prestes a tomar as maiores capitais da Terra. Justamente quando estão esmagando todas as resistências, reduzindo nossa civilização a escombros, eles interrompem misteriosamente a investida. Apesar de todo o avanço de sua ciência e armamentos, eles não levaram em conta o ataque da mais simplória das criaturas: os germes.

Esse romance deu origem a todo um gênero que lançou mil filmes, desde *A invasão dos discos voadores* até *Independence Day*. Muitos cientistas, porém, estremecem ao ver a caracterização dos alienígenas como criaturas dotadas de valores e emoções humanas. Apesar da pele verde lustrosa e da cabeça enorme, até certo ponto eles ainda se parecem conosco. E geralmente falam inglês com perfeição.

Mas, na opinião de muitos cientistas, nós podemos ter muito mais em comum com uma lagosta ou um caramujo do que com um alienígena vindo do espaço.

Tal como a consciência de silício, é muito mais provável que a consciência alienígena tenha as características gerais descritas em nossa teoria do espaço-tempo, isto é, a capacidade de criar um modelo do mundo e calcular como irá evoluir no tempo a fim de atingir metas. Mas, enquanto os robôs são programados para ter laços emocionais com os humanos e metas compatíveis com as nossas, a consciência alienígena pode não ter nada disso. É mais provável que tenham seus próprios valores e metas, independentes da humanidade. Só nos resta especular quais seriam.

O físico dr. Freeman Dyson, do Institute for Advanced Study em Princeton, foi consultor no filme *2001: Uma odisséia no espaço*. Quando assistiu ao filme ficou encantado, não com os espetaculares efeitos especiais, mas porque era o primeiro filme de Hollywood a apresentar uma consciência alienígena com desejos, metas e intenções totalmente diferentes dos nossos. Pela primeira vez, os alienígenas não eram atores humanos circulando com fantasias de monstros cafonas, fingindo ser criaturas ameaçadoras. Em vez disso, foi apresentada uma consciência alienígena totalmente ortogonal à experiência humana, algo inteiramente além do nosso conhecimento.

Em 2011, Stephen Hawking levantou outra questão. O conceituado cosmólogo ganhou as manchetes ao dizer que devemos estar preparados para um possível ataque alienígena. Ele disse que, se algum dia viermos a encontrar outra civilização, será mais avançada que a nossa e, portanto, um perigo mortal para nossa existência.

Basta lembrar o que aconteceu com os astecas na chegada do sanguinário

Hernán Cortés e seu séquito de conquistadores para imaginar o que nos aconteceria num malfadado encontro desses. Armados com uma tecnologia que os astecas, da Idade do Bronze, nunca tinham visto, como espadas de ferro, pólvora e cavalos, aquele pequeno bando de degoladores esmagou a antiga civilização asteca em questão de meses, em 1521.

Tudo isso traz outras questões: Como será a consciência alienígena? Como seus processos de pensamento e metas diferem dos nossos? O que eles querem?

PRIMEIRO CONTATO NESTE SÉCULO

Não se trata de uma questão acadêmica. Em vista dos notáveis avanços na astrofísica, é bem possível que realmente consigamos fazer contato com uma inteligência alienígena nas próximas décadas. Nossa reação a tal contato pode determinar um dos eventos mais cruciais da história da humanidade.

Vários avanços têm tornado esse dia possível.

Primeiro, em 2011, o satélite Kepler, pela primeira vez na história, deu aos cientistas um “censo” da Via Láctea. Após analisar a luz de milhares de estrelas, o Kepler indicou que uma em cada duzentas estrelas poderia abrigar um planeta semelhante à Terra, numa zona habitável. Pela primeira vez, podemos calcular quantas estrelas semelhantes à da Terra podem existir na Via Láctea: cerca de um bilhão. Quando olhamos para as estrelas no céu, temos razão em presumir que alguém lá está olhando para nós.

Até o momento, mais de mil exoplanetas foram analisados minuciosamente por telescópios situados na Terra (os astrônomos descobrem uma média de dois exoplanetas por semana). Infelizmente, quase todos são do tamanho de Júpiter, provavelmente desprovidos de criaturas parecidas com as nossas, mas há alguns planetas rochosos só algumas vezes maiores que a Terra, as “superTerras”. O satélite Kepler já identificou cerca de 2.500 possíveis exoplanetas, alguns deles muito parecidos com a Terra. Tais planetas estão a uma distância da estrela mãe propícia à existência de oceanos líquidos. E a água líquida é o “solvente universal”, que dissolve a maioria das substâncias químicas orgânicas, como o DNA e as proteínas.

Em 2013, cientistas da Nasa anunciaram a descoberta mais espetacular com o Kepler: dois exoplanetas quase gêmeos da Terra. Estão a 1.200 anos-luz daqui, na constelação de Lyra. São apenas 60% e 40% maiores que a Terra e, o mais importante, ambos estão na zona habitável da estrela mãe, indicando a possibilidade de terem oceanos líquidos. Entre todos os planetas analisados até agora, esses estão mais perto de espelhar a Terra.

Além disso, o telescópio espacial Hubble nos deu uma estimativa do número

de galáxias no universo visível: cem bilhões. Desse modo, podemos calcular que o número de planetas semelhantes à Terra no universo visível é de um bilhão vezes cem bilhões, ou seja, há cem quintilhões de planetas parecidos com a Terra.

É realmente um número astronômico. Portanto, as chances de existir vida no universo são astronômicas, principalmente quando se considera que ele tem 13,8 bilhões de anos de idade, tempo suficiente para a ascensão – e queda – de muitos impérios inteligentes. Na verdade, o milagre seria se *não* existissem outras civilizações avançadas.

PROJETO SETI E CIVILIZAÇÕES ALIENÍGENAS

A tecnologia de radiotelescópio está ficando mais sofisticada. Até o momento, apenas cerca de mil estrelas foram analisadas detalhadamente na busca de sinais de vida inteligente, mas na próxima década esse número deve aumentar um milhão de vezes.

A procura de civilizações alienígenas por esse meio data de 1960, quando o astrônomo Frank Drake deu início ao Project Ozma (numa referência à Rainha de Oz), com um radiotelescópio de 25 metros em Green Bank, na Virgínia Ocidental. Isso marcou o nascimento do projeto Seti (Search for Extraterrestrial Intelligence). Infelizmente, não foram captados sinais de alienígenas, mas em 1971 a Nasa propôs o Projeto Cyclops, que teria 1.500 radiotelescópios, a um custo de dez bilhões de dólares.

Como era de esperar, não deu em nada. O Congresso não achou a menor graça.

Conseguiram verba para um projeto, bem mais modesto, de enviar uma mensagem cuidadosamente codificada para alienígenas no espaço sideral, em 1971. A mensagem codificada, contendo 1.679 bits de informação, foi transmitida através do gigantesco radiotelescópio de Arecibo, em Porto Rico, direcionada ao Grande Aglomerado Globular M13, a cerca de 25.100 anos-luz de distância. Foi o primeiro cartão de visita cósmico da Terra, contendo informações relevantes sobre a raça humana. Até hoje não teve resposta. Talvez os alienígenas não tenham ficado impressionados conosco, ou, quem sabe, a velocidade da luz atrapalhou. Dada a grande distância, a data mais próxima para receber uma resposta é daqui a 52.174 anos.

Desde então, alguns cientistas manifestam receio de anunciar nossa existência a alienígenas, pelo menos até sabermos suas intenções para conosco. E discordam dos proponentes do projeto Meti (Messaging to Extraterrestrial Intelligence), que promovem ativamente o envio de sinais a civilizações

alienígenas. A justificativa do projeto Meti é que, se a Terra já manda grandes quantidades de sinais de rádio e TV para o espaço sideral, algumas mensagens a mais não farão muita diferença. Mas os críticos do Meti acham que não devemos aumentar as chances de sermos descobertos por alienígenas potencialmente hostis.

Em 1995, astrônomos recorreram à iniciativa privada para instalar o Instituto Seti em Mountain View, na Califórnia, a fim de centralizar as pesquisas e dar início ao Projeto Phoenix, que estuda mil estrelas próximas semelhantes ao Sol na faixa de rádio entre 1.200 e 3.000 megahertz. O equipamento é tão sensível que pode captar emissões do sistema de radar de um aeroporto a duzentos anos-luz de distância. Desde a sua fundação, o Instituto Seti já fez varreduras em mais de mil estrelas, a um custo de cinco milhões de dólares por ano, mas ainda não teve sucesso.

Uma abordagem mais recente é o projeto SETI@home, lançado em 1999 por astrônomos da Universidade da Califórnia em Berkeley, que conta com um exército informal de milhões de amadores, proprietários de microcomputadores. Qualquer um pode entrar nessa caçada histórica. À noite, enquanto o voluntário dorme, o protetor de tela do computador processa alguns dados que fluem pelo radiotelescópio de Arecibo em Porto Rico. Até o momento, o projeto tem 5,2 milhões de usuários, em 234 países. Talvez esses amadores sonhem em ser os primeiros humanos a fazer contato com extraterrestres. O nome do felizardo entraria para a história, como o de Colombo. O SETI@home cresceu tão rapidamente que hoje é o maior empreendimento computacional desse tipo.

Quando entrevistei o dr. Dan Wertheimer, diretor do SETI@home, perguntei como ele conseguia distinguir entre mensagens falsas e reais. Sua resposta me surpreendeu. Ele disse que às vezes “plantam” deliberadamente nos dados do radiotelescópio sinais falsos de uma civilização inteligente fictícia. Se ninguém capta essas mensagens falsas, eles sabem que há algo de errado com o sistema. Aprendam a lição: se a tela do seu computador anunciar que decifrou uma mensagem de uma civilização extraterrestre, não telefone imediatamente para a polícia nem para o presidente da República, porque pode ser uma mensagem falsa.

CAÇADORES DE ALIENÍGENAS

Um colega meu dedicado a encontrar vida inteligente no espaço é o dr. Seth Shostak, diretor do Instituto Seti. Sendo ele Ph.D. em física pelo California Institute of Technology, achei que fosse se tornar um ilustre professor de física, dando aulas para doutorandos ávidos de saber, mas não. Ele se dedica a funções

muito diferentes: pede a pessoas ricas doações para o Instituto Seti, passa horas à espera de sinais de extraterrestres, e tem um programa de rádio. Certa vez, perguntei a ele sobre o “fator risadinha” – nossos colegas cientistas dão uma risadinha quando ele conta que fica à escuta de vozes do espaço sideral? Não mais, ele afirma. Com todas as novas descobertas da astronomia, a maré virou.

Na verdade, ele não se intimida e diz, tranquilamente, que fará contato com uma civilização alienígena num futuro muito próximo. E declarou publicamente que o Allen Telescope Array de 350 antenas que está sendo construído agora, vai “cruzar com um sinal até o ano 2025”.

Não é um pouco arriscado?, perguntei. Como ele tem tanta certeza? Um ponto a seu favor é a explosão do número de radiotelescópios nos últimos anos. Embora o governo dos Estados Unidos não dê apoio financeiro ao projeto, o Instituto Seti recentemente achou uma mina de ouro quando convenceu Paul Allen (o bilionário da Microsoft) a doar mais de trinta milhões de dólares para criar o Allen Telescope Array em Hat Creek, na Califórnia, 467 quilômetros ao norte de San Francisco. Atualmente são 42 radiotelescópios, mas serão 350 varrendo o céu. (Um problema é a falta crônica de financiamento para esses experimentos científicos. Para compensar o orçamento apertado, a instalação em Hat Creek se mantém graças a um financiamento parcial das Forças Armadas.)

Shostak confessa que se incomoda um pouco quando as pessoas confundem o projeto Seti com caçadores de óvnis. O Seti, ele afirma, se fundamenta em princípios sólidos da física e da astronomia, usando tecnologia de ponta. Os caçadores de óvnis baseiam suas teorias em contadores de casos, em provas indiretas que podem ou não se fundamentar na verdade. O problema é que a enxurrada de relatos de aparições de óvnis que ele recebe por correio não é reproduzível nem testável. Ele pede a todos que forem abduzidos por um disco voador que roubem alguma coisa – um peso de papel ou caneta alienígena, por exemplo – como prova da história. Nunca saia de um disco voador de mãos vazias, ele me aconselhou.

E conclui que não há provas concretas de visitas de alienígenas a nosso planeta. Perguntei se ele achava que o governo dos Estados Unidos escondia as evidências desses encontros, como sugere a teoria da conspiração. Ele respondeu: “E eles teriam competência para esconder algo desse tamanho? Lembre-se de que é esse mesmo governo que gerencia os correios.”

A EQUAÇÃO DE DRAKE

Quando perguntei ao dr. Wertheimer por que ele tinha tanta certeza da existência de vida inteligente no espaço, ele respondeu que os números estavam a seu favor.

Em 1961, o astrônomo Frank Drake fez suposições plausíveis para estimar a quantidade de civilizações inteligentes. Se começarmos com cem bilhões, o número de estrelas na Via Láctea, podemos estimar a fração delas que são semelhantes ao nosso Sol. E podemos reduzir esse número estimando a fração delas que tem planetas, quantos deles são semelhantes à Terra etc. Chegando a um número razoável de suposições, temos uma estimativa de dez mil civilizações avançadas na Via Láctea. (Carl Sagan, com outro conjunto de estimativas, chegou ao número de um milhão.)

Desde então, os cientistas puderam fazer estimativas muito melhores do número de civilizações avançadas em nossa galáxia. Por exemplo: sabemos que existem mais planetas orbitando estrelas do que Drake esperava, e também mais planetas semelhantes à Terra. Mas ainda temos um problema. Mesmo se soubéssemos quantos planetas iguais ao nosso existem no espaço, ainda não saberíamos quantos deles têm vida inteligente. Na Terra, levou 4,5 bilhões de anos até que seres inteligentes (nós) finalmente saíssem dos pântanos. Já existiam formas de vida aqui há 3,5 bilhões de anos, mas apenas nos últimos cem mil anos, mais ou menos, surgiram seres inteligentes como nós. Portanto, mesmo num planeta igual à Terra, o surgimento de vida inteligente é muito difícil.

POR QUE ELES NÃO VÊM AQUI?

Depois fiz ao dr. Seth Shostak, do Seti, a pergunta fatal: Se existem tantas estrelas na galáxia, e tantas civilizações alienígenas, por que eles não vêm aqui? Esse é o paradoxo de Fermi, assim batizado por causa de Enrico Fermi, ganhador do Prêmio Nobel, que ajudou a construir a bomba atômica e descobriu os segredos do núcleo do átomo.

Muitas teorias foram propostas. Uma diz que a distância entre estrelas deve ser grande demais. Levaria cerca de 70 mil anos para nossos foguetes químicos mais potentes atingirem as estrelas mais próximas da Terra. Talvez uma civilização milhares ou milhões de anos mais avançada que a nossa solucione esse problema, mas há outra possibilidade. Talvez tenham se aniquilado numa guerra nuclear. Como John F. Kennedy falou: “Lamento dizer, mas talvez faça sentido supor que a vida se extinguiu em outros planetas porque os cientistas de lá eram mais avançados que os nossos.”

Ou talvez a razão mais lógica seja a seguinte: Imagine que estamos andando pelo campo e encontramos um formigueiro. Vamos nos abaixar e dizer para as formigas: “Vou trazer colares para vocês, vou trazer espelinhos, vou lhes dar energia nuclear. Vou criar um paraíso das formigas. Quero falar com seu líder.”?

Provavelmente, não.

Agora imagine operários construindo uma supervia expressa com oito pistas, ao lado do formigueiro. As formigas vão saber em qual frequência os operários estão falando? Vão saber o que é uma supervia expressa com oito pistas? Da mesma forma, qualquer civilização inteligente que desça à Terra estará milhares, ou milhões, de anos à nossa frente, e não teremos nada a oferecer a eles. Em outras palavras, é arrogância nossa acreditar que alienígenas vão viajar trilhões e trilhões de quilômetros só para nos ver.

Muito provavelmente, não estamos na tela de radar deles. Ironicamente, a galáxia pode estar cheia de formas de vida inteligente e somos tão primitivos que nem as notamos.

PRIMEIRO CONTATO

Mas suponhamos que virá o momento, talvez mais cedo do que se pensa, em que faremos contato com uma civilização alienígena. Isso pode significar um ponto de virada na história da humanidade. Portanto, as próximas questões são: o que eles querem, e como será a consciência deles?

Nos filmes e histórias de ficção científica, os alienígenas geralmente querem nos devorar, conquistar, acasalar conosco, escravizar, ou arrancar todos os recursos valiosos do nosso planeta. Tudo isso é altamente improvável.

Possivelmente, nosso primeiro contato com uma civilização alienígena não será com um disco voador aterrissando no gramado da Casa Branca. É mais provável que um adolescente do projeto SETI@home venha dizer que seu computador decodificou sinais do radiotelescópio de Arecibo em Porto Rico. Ou que o Seti de Hat Creek detecte uma mensagem indicando inteligência no espaço.

Nosso primeiro encontro será um evento de mão única. Podemos ficar na escuta e responder a uma mensagem, mas essa resposta levará décadas ou séculos para chegar lá.

O que ouvirmos pelo rádio poderá nos dar pistas valiosas dessa civilização alienígena, mas a maioria das mensagens será provavelmente de variedades, fofocas, músicas etc., e muito pouco conteúdo científico.

Fiz ao dr. Shostak a segunda pergunta fatal: Você vai manter segredo quando for feito o Primeiro Contato? Afinal, isso não irá causar pânico em massa, histeria religiosa, caos, fuga espontânea das cidades? Fiquei meio surpreso quando ele respondeu que não. Todos os dados seriam entregues a todos os governos e a todos os povos do mundo.

As próximas questões são: Como eles são? O que eles pensam?

Para entender a consciência alienígena, talvez seja instrutivo analisar outra consciência muito diferente da nossa: a consciência dos animais. Vivemos com

eles, mas ignoramos totalmente o que se passa na mente deles.

E entender a consciência animal pode nos ajudar a entender a consciência alienígena.

CONSCIÊNCIA ANIMAL

Os animais pensam? Se pensam, sobre o que pensam? Há milhares de anos que essa pergunta deixa grandes filósofos sem resposta. Os historiadores gregos Plutarco e Plínio escreveram sobre uma famosa questão que até hoje não foi resolvida. No decorrer dos séculos, diversas soluções foram apresentadas por gigantes da filosofia.

Um cachorro vem andando por uma estrada, procurando seu dono, e chega a uma encruzilhada. O cachorro pega o caminho da esquerda, fareja, e retorna, sabendo que o dono não passou por ali. Depois o cachorro pega o caminho da direita, fareja, e volta, sabendo que o dono não seguiu por ali também. Então o cachorro, decididamente, toma o caminho do meio, sem farejar.

O que se passou na mente do cachorro? Grandes pensadores se debruçaram sobre essa questão, sem sucesso. O filósofo e ensaísta francês Michel de Montaigne escreveu que o cachorro concluiu obviamente que a única solução possível era seguir o caminho do meio, demonstrando que os cachorros são capazes de ter pensamento abstrato.

No século XIII, São Tomás de Aquino dizia o oposto – que aparentar ter pensamento abstrato não é o mesmo que realmente tê-lo. A aparência superficial de inteligência pode nos enganar, ele dizia.

Séculos depois houve um famoso debate entre John Locke e George Berkeley sobre consciência animal. “Bichos não abstraem”, disse Locke, ao que o bispo Berkeley respondeu: “Se o fato de que as bestas não abstraem é a propriedade que as distingue, receio que muitos dos que se passam por homens devam ser incorporados à espécie delas.”

Outros filósofos, em várias épocas, tentaram analisar a questão da mesma maneira, isto é, impondo a consciência humana ao cachorro. É um equívoco do antropomorfismo supor que os animais pensam e se comportam como nós. E talvez a solução seja olhar essa questão do ponto de vista do cachorro, que poderia ser bem alienígena.

Na definição de consciência dada no capítulo 2, os animais são parte de um *continuum* de consciência. Diferem de nós quanto aos parâmetros que usam para criar um modelo do mundo. O dr. David Eagleman diz que os psicólogos chamam a isso de *umwelt*, isto é, a realidade percebida por outros animais. Ele diz: “No mundo surdo e cego do carrapato, os sinais importantes são a

temperatura e o cheiro do ácido butírico. Para um peixe elétrico, são os campos elétricos. Para o morcego, que usa o eco como meio de localização, são as ondas de ar comprimido. Cada organismo tem sua própria *umwelt*, e possivelmente supõe ser a realidade objetiva total 'lá de fora'."

Consideremos o cérebro de um cachorro, que vive todo o tempo num turbilhão de odores e assim encontra comida e parceira para acasalar. A partir desses odores, o cão constrói um mapa mental do que existe ao seu redor. Esse mapa de odores é totalmente diferente do que construímos com os olhos, e transmite um conjunto de informações totalmente diferente. (Lembre do dr. Penfield, no capítulo 1, que desenhou um mapa do córtex cerebral com a imagem distorcida da correspondência com os órgãos do corpo [o Homenzinho de Penfield]. Agora imagine um diagrama do cérebro de um cão desenhado por Penfield. A parte maior seria do focinho, e não dos dedos. Cada animal teria um diagrama de Penfield totalmente diferente. Os alienígenas teriam um diagrama de Penfield ainda mais estranho.)

Infelizmente, temos tendência a atribuir uma consciência humana aos animais, embora eles tenham uma perspectiva do mundo totalmente diferente. Por exemplo: ao ver um cachorro seguir fielmente as ordens do dono, supomos, inconscientemente, que o cão é o melhor amigo do homem porque gosta de nós e nos respeita. Mas como o cachorro é descendente do *Canis lupus*, o lobo cinzento, que caça em grupo dentro de uma hierarquia rígida, é mais provável que o cachorro nos veja como um tipo de macho alfa, líder da matilha. Em certo sentido, você é o Cachorro Chefe. (Provavelmente por isso é muito mais fácil treinar filhotinhos do que cães velhos. É mais fácil marcar sua presença no cérebro do filhote, ao passo que os cães adultos já percebem que os humanos não fazem parte do bando.)

Além disso, quando um gato entra na sala e urina no tapete, achamos que o gato está zangado ou nervoso, e tentamos encontrar o motivo para o aborrecimento do gato. Mas talvez o gato esteja simplesmente marcando território com o cheiro da urina para espantar outros gatos. Então o gato não está nada aborrecido, mas só avisando a outros gatos que fiquem longe da casa porque a casa pertence a ele.

E se o gato ronrona e se esfrega nas suas pernas, você acha que ele está demonstrando gratidão por você cuidar dele, que se trata de um sinal de carinho e afeição. É muito mais provável que o gato esteja esfregando os hormônios dele nas suas pernas para marcar propriedade (no caso, você) e repelir outros gatos. Do ponto de vista do gato, você é uma espécie de servo treinado para lhe dar comida várias vezes por dia, e a esfregação do cheiro é para avisar a outros gatos que não se aproximem do servo dele.

Como disse Michel de Montaigne, o grande filósofo do século XVI: "Quando brinco com minha gata, como saber se é ela que está brincando comigo, e não eu

com ela?”

E se o gato se retira para ficar sozinho, não é sinal de raiva nem de indiferença. Ele é descendente do gato selvagem, que é um caçador solitário, ao contrário do cachorro. Não fica babando à disposição de um macho alfa, como é o caso dos cães. A proliferação de programas de “encantadores de animais” na televisão é um sinal dos problemas encontrados quando atribuímos aos animais a consciência e as intenções humanas.

Um morcego também deve ter uma consciência muito diferente, dominada por sons. Quase cego, o morcego precisa de feedback contínuo dos guinchos que emite para localizar insetos, obstáculos e outros morcegos via sonar. O mapa de Penfield do cérebro de um morcego seria muito diferente do nosso, com uma parte enorme tomada pelas orelhas. Assim como o dos golfinhos, que têm uma consciência diferente da dos humanos, muito baseada no sonar. Como os golfinhos têm um córtex frontal menor, pensava-se que não eram inteligentes, mas isso é compensado por uma massa cerebral maior. O desdobramento do neocórtex do golfinho cobre seis páginas de revista, ao passo que o de um ser humano cobre quatro páginas. O golfinho tem os córtices parietais e temporais muito desenvolvidos para analisar sinais de sonar na água, e é um dos poucos animais capazes de se reconhecer no espelho, provavelmente devido a esse desenvolvimento cortical.

Além disso, o cérebro do golfinho tem uma estrutura diferente da dos humanos porque a linhagem de golfinhos e humanos divergiu há cerca de 95 milhões de anos. Como o golfinho não tem necessidade de ter um nariz, seu bulbo olfativo desaparece logo após o nascimento. Mas há 30 milhões de anos o córtex auditivo deles atingiu um tamanho enorme porque aprenderam a usar a ecolocalização, o sonar, para encontrar comida. O mundo deles, assim como o dos morcegos, deve ser um turbilhão de ecos e vibrações. Em comparação com os humanos, os golfinhos têm um lobo extra no sistema límbico, chamado região paralímbica, que provavelmente os ajuda a formar fortes relações sociais.

Os golfinhos têm uma linguagem inteligente. Certa vez entrei numa piscina com golfinhos para a gravação de um documentário do Science Channel. Coloquei sonares na piscina para captar os cliques e assobios que eles usam para se comunicar. Os sinais foram captados e analisados por computador. Existe um modo simples de discernir se há uma inteligência oculta naqueles conjuntos de guinchos e trinados. Na língua inglesa, por exemplo, a letra *e* é a mais usada. De fato, podemos criar uma lista da frequência de uso de todas as letras do alfabeto. Qualquer livro em inglês que analisemos por computador irá mostrar praticamente a mesma lista de letras mais encontradas.

Da mesma forma, esse programa de computador pode ser usado para analisar a linguagem dos golfinhos, e encontramos um padrão indicando inteligência. Entretanto, em outros mamíferos, o padrão começa a falhar e termina se

quebrando totalmente à medida que chegamos nos animais inferiores, com cérebro menor. Nesse caso, o sinal é quase aleatório.

ABELHAS INTELIGENTES?

Para uma noção de como deve ser a consciência alienígena, consideremos a estratégia adotada pela natureza para reproduzir a vida na Terra. São duas estratégias reprodutivas básicas, com profundas implicações na evolução e na consciência.

A primeira, usada pelos mamíferos, é produzir um pequeno número de rebentos e cuidar de cada um até chegar à maturidade. É uma estratégia arriscada, porque poucos indivíduos são produzidos a cada geração, e dependem da criação para enfrentar as adversidades. Isso significa que cada vida é bem acolhida e bem cuidada por um certo período de tempo.

Outra estratégia, muito mais antiga, é usada pela maior parte do reino vegetal e animal, incluindo os insetos, répteis e quase todas as outras formas de vida do planeta. Trata-se de gerar um grande número de ovos ou sementes, e deixá-los germinar por si mesmos. Sem os cuidados da criação, a maioria não sobrevive e apenas uns poucos indivíduos mais resistentes produzem uma nova geração. Isso significa que a energia investida pelos pais em cada geração é nula e a reprodução depende da lei das médias para propagar a espécie.

Essas duas estratégias resultam em atitudes muito diferentes com relação à vida e à inteligência. A primeira valoriza cada indivíduo. Amor, cuidados e laço afetivo têm um valor determinante nesse grupo. Essa estratégia reprodutiva só dá bom resultado quando os pais investem uma quantidade considerável de energia para preservar as crias. A segunda estratégia, por sua vez, não valoriza o indivíduo, e enfatiza a sobrevivência da espécie ou do grupo como um todo. Para eles, a individualidade não tem significado.

A estratégia reprodutiva tem implicações profundas na evolução da inteligência. Quando duas formigas se encontram, por exemplo, trocam um número limitado de informações por meio de odores químicos e gestos. Apesar de mínima, a informação trocada é capaz de levá-las a criar túneis e galerias elaboradas para construir o formigueiro. Da mesma forma, as abelhas, embora se comuniquem por meio de uma dança, conseguem fazer colmeias complexas e localizar canteiros de flores muito distantes trabalhando de forma coletiva. A inteligência delas não é formada individualmente, mas vem da interação holística de toda a colônia e dos seus genes.

Portanto, imagine uma civilização extraterrestre inteligente baseada na segunda estratégia, como uma raça de abelhas inteligentes. Nessa sociedade, as

abelhas operárias que voam todo dia à procura de pólen são descartáveis. As operárias não se reproduzem e vivem com o único propósito de servir à colmeia e à rainha, por quem elas se sacrificam sem hesitação. Os vínculos entre os mamíferos não têm significado para elas.

Hipoteticamente, isso poderia afetar o programa espacial delas. Como nós valorizamos muito a vida dos astronautas, empregamos grandes recursos para trazê-los de volta vivos. Boa parte do custo das viagens espaciais é aplicada na garantia da saúde da tripulação, na viagem de volta e na reentrada na atmosfera. Mas numa civilização de abelhas inteligentes a vida de cada operária não vale tanto, por isso o programa espacial delas tem um custo muito baixo. As operárias não precisam voltar; cada viagem pode ser só de ida, o que representa uma economia significativa.

Agora, imagine se encontramos um extraterrestre com função similar à de uma abelha operária. Normalmente, se encontramos uma abelha na floresta, a probabilidade maior é de que ela nos ignore, a não ser que ela ou a colmeia se sintam ameaçadas. É como se não existíssemos. Da mesma forma, é provável que esse operário espacial não tenha o menor interesse em fazer contato conosco, e muito menos em compartilhar seu saber. Ele seguiria cumprindo sua missão e nos ignoraria. E os valores que prezamos não significariam nada para ele.

Nos anos 1970, foram colocadas placas nas sondas *Pioneer 10* e *11*, contendo informações sobre nosso mundo e nossa sociedade, exaltando a diversidade e a riqueza da vida na Terra. Na época, os cientistas supunham que civilizações extraterrestres seriam como nós, curiosas e interessadas em fazer contato. Mas se um alienígena operário como as abelhas encontrasse nossas placas, é provável que não significassem nada para ele.

Além disso, os operários não precisam ser muito inteligentes. Só precisam ser inteligentes o suficiente para servir aos interesses do grupo. Portanto, se enviamos uma mensagem a um planeta de abelhas inteligentes, o mais provável é que elas não tenham interesse em mandar uma mensagem de volta.

Ainda que seja possível fazer contato com uma civilização dessas, a comunicação deve ser muito difícil. Por exemplo: quando nos comunicamos uns com os outros, dividimos as ideias em frases com uma estrutura de sujeito e verbo, a fim de construir uma narrativa, geralmente uma história pessoal. A maioria de nossas frases tem a seguinte estrutura: “Eu fiz isso”, ou “Eles fizeram aquilo”. A maior parte de nossa literatura e conversas é composta por narrativas, relatos envolvendo experiências e aventuras que nós ou nossos personagens tivemos. Isso pressupõe que nossas experiências pessoais são a forma dominante de transmitir informação.

No entanto, uma civilização baseada em abelhas inteligentes pode não ter o menor interesse em histórias e experiências subjetivas. Sendo uma civilização

altamente coletivizada, as mensagens não devem ser pessoais, mas puramente factuais, contendo informações necessárias à vida da colmeia, e não amenidades e mexericos para elevar a posição social individual. Devem até achar nossas histórias meio repulsivas, pois colocam o papel do indivíduo na frente das necessidades coletivas.

Além disso, abelhas operárias devem ter uma noção de tempo totalmente diferente. Como elas são descartáveis, podem não ter um tempo de vida muito longo. Seus projetos devem ser de curto prazo e bem definidos.

No entanto, nós, humanos, vivemos muito mais; mas também temos uma noção de tempo tácita. Nossos projetos e ocupações têm um prazo de conclusão enquanto estamos vivos. De modo inconsciente, traçamos nossos projetos, nossas relações com os outros e nossos objetivos de modo a serem cumpridos dentro do nosso período de vida. Em outras palavras, temos fases distintas na vida: solteiro, casado, criando filhos, aposentado. Talvez sem termos plena consciência disso o tempo todo, sabemos que vamos viver um determinado tempo e depois morrer, que nossa vida é finita.

Mas imagine seres que podem viver milhares de anos, ou que são imortais. Suas prioridades, objetivos e ambições têm que ser completamente diferentes. Podem ter projetos que levariam muitas gerações de vida humana. A viagem interestelar é considerada pura ficção científica porque, como vimos, um foguete convencional leva cerca de 70 mil anos para chegar às estrelas mais próximas. Para nós, é uma demora proibitiva. Para uma forma de vida alienígena, porém, esse tempo pode ser irrelevante. Por exemplo: eles podem ser capazes de hibernar, desacelerar o metabolismo, ou até mesmo de viver indefinidamente.

COMO ELES SÃO?

As primeiras traduções de mensagens alienígenas provavelmente nos darão pistas de sua cultura e modo de vida. Por exemplo: é possível que tenham evoluído de predadores e, conseqüentemente, tenham conservado algumas características deles. (Em geral, os predadores terráqueos são mais espertos que as presas. Caçadores, como tigres, leões, gatos e cães, usam a astúcia para perseguir, se esconder e tocar, e tudo isso requer inteligência. Todos esses predadores têm olhos na frente da cara, para uma visão estéreo enquanto focalizam a atenção. As presas, que têm olhos nos lados da cara, ao verem o predador só lhes resta fugir. Por isso é que se diz “dissimulado como uma raposa” ou “um coelhinho assustado”.) As formas de vida alienígenas podem ter deixado para trás muitos instintos predadores de seus ancestrais, mas é possível que ainda tenham um pouco da consciência deles (isto é, territorialidade, expansão e

violência quando for necessário).

Ao examinar a raça humana, identificamos pelo menos três elementos básicos que se conjugaram para nos tornar inteligentes:

1. Polegar oponível, que nos confere a habilidade de manipular e modelar o ambiente utilizando ferramentas.
2. Olhos estéreo ou os olhos 3D de um caçador.
3. Linguagem, que nos permite acumular conhecimento, cultura e sabedoria com o passar das gerações.

Quando comparamos esses três elementos com características encontradas no reino animal, percebemos que são poucos os animais que preenchem esses requisitos de inteligência. Gatos e cachorros, por exemplo, não têm uma linguagem complexa, nem a habilidade de pegar um objeto. Polvos possuem tentáculos sofisticados, mas enxergam mal e não têm uma linguagem complexa.

Pode ser que haja adaptações nesses critérios. Em vez de polegares oponíveis, um alienígena pode ter garras ou tentáculos (o único pré-requisito é que possam manipular o ambiente com instrumentos criados por esses apêndices). Em vez de dois olhos, eles podem ter muitos, como os insetos, ou sensores que detectem sons ou radiação ultravioleta em vez da luz visível. É mais provável que tenham a visão estéreo de um caçador, porque os predadores em geral têm inteligência superior à da presa. E em vez da linguagem comunicada por sons, talvez se comuniquem por diversas formas de vibrações (a única exigência é que troquem informações entre si para criar uma cultura que se estenda por várias gerações).

Fora esses três critérios, vale tudo.

Os alienígenas devem ter uma consciência matizada pelo ambiente. Hoje os astrônomos entendem que talvez o habitat mais rico no universo não sejam os planetas semelhantes à Terra, que oferecem o deleite da luz cálida da estrela mãe, e sim satélites gelados na órbita de planetas grandes como Júpiter, a bilhões de quilômetros da estrela. Acredita-se que Europa, uma lua gelada de Júpiter, tenha um oceano líquido por baixo do gelo da superfície, aquecido por forças de maré. Como Europa tem órbita excêntrica, é comprimida em várias direções pela imensa força gravitacional de Júpiter, o que causa fricção dentro da lua. Isso gera calor, formando vulcões e fontes hidrotermais que derretem o gelo e criam oceanos líquidos. Estima-se que os oceanos de Europa sejam muito profundos e seu volume pode ser muito maior que o dos oceanos da Terra. Sabendo-se que 50% das estrelas no céu devem ter planetas do tamanho de Júpiter (cem vezes mais abundantes que os semelhantes à Terra), as formas de vida mais ricas devem estar nas luas geladas de gigantes gasosos como Júpiter.

Sendo assim, a primeira civilização alienígena que encontrarmos no espaço muito provavelmente será formada por seres de origem aquática. É provável também que eles tenham migrado dos oceanos e aprendido a viver na superfície gelada, longe da água líquida, por várias razões. Primeiro, qualquer espécie que vive apenas sob o gelo terá uma visão muito limitada do universo. Não poderão ter desenvolvido a astronomia nem programas espaciais se pensarem que o universo é apenas um oceano sob uma capa de gelo. Segundo, componentes elétricos entram em curto-circuito em contato com a água, logo, eles não poderão ter desenvolvido rádio nem televisão se viverem debaixo d'água. Para essa civilização avançar, precisa dominar a eletrônica, que não pode existir nos mares. Portanto, o mais provável é que esses alienígenas tenham deixado o oceano e aprendido a sobreviver na superfície, como o nós fizemos.

Mas, e se essa forma de vida evoluiu para uma civilização capaz de criar um programa espacial e chegar à Terra? Ainda serão organismos biológicos como nós, ou já serão pós-biológicos?

A ERA PÓS-BIOLÓGICA

Um dos cientistas que tem se dedicado a essas questões é meu colega dr. Paul Davies, da Universidade do Estado do Arizona perto de Phoenix. Quando o entrevistei, ele me falou que precisamos expandir nossos horizontes para contemplar a aparência que poderá ter uma civilização mil anos ou mais à nossa frente.

Em vista dos perigos de viagens espaciais, ele acredita que tais seres terão abandonado a forma biológica, a exemplo das mentes incorpóreas que vimos no capítulo anterior. Ele escreve: “Minha conclusão é espantosa. Penso ser muito provável – na verdade, inevitável – que a inteligência biológica seja apenas um fenômeno transitório, um estágio fugaz na evolução da inteligência no universo. Se algum dia viermos a encontrar uma inteligência extraterrestre, há uma probabilidade esmagadora de que seja de natureza pós-biológica, uma conclusão que tem ramificações óbvias e de longo alcance para o Seti.”

De fato, se os alienígenas estão milhares de anos à nossa frente, é muito provável que tenham abandonado o corpo biológico há muitas eras para criar o mais eficiente corpo computacional: um planeta com a superfície totalmente coberta de computadores. Davies diz: “Não é difícil imaginar a superfície inteira de um planeta coberta por um único sistema de processamento integrado. (...) Ray Bradbury cunhou o termo ‘cérebros Matrioshka’ para definir essas entidades impressionantes.”

Para Davies, a consciência alienígena deve perder o conceito de “si mesmo” e

ser absorvida pela “World Wide Web de Mentes”, que cobre toda a superfície do planeta. Ele acrescenta que “Uma rede de computadores de grande potência sem o senso de si mesmo terá uma enorme vantagem sobre a inteligência humana porque pode reestruturar a ‘si mesma’, fazer mudanças sem temor, fundir-se com outros sistemas, e crescer. ‘Sentimentos pessoais’ são um claro impedimento ao progresso”.

Portanto, em nome da eficiência e maior capacidade computacional, ele imagina os integrantes dessa civilização avançada abrindo mão de sua identidade e sendo absorvidos por uma consciência coletiva.

Davies reconhece que os críticos dessa ideia devem achar esse conceito totalmente repulsivo. Parece que essa espécie alienígena sacrifica a individualidade e a criatividade a favor do bem maior do coletivo ou da colmeia. Ele adverte que isso não é inevitável, mas é a opção mais eficiente para a civilização.

Davies tem uma hipótese que ele mesmo admite ser muito deprimente. Quando perguntei por que essas civilizações não nos visitam, ele me deu uma resposta estranha. Disse que uma civilização nesse estágio tão avançado terá desenvolvido também realidades virtuais muito mais interessantes e desafiadoras que a realidade. A realidade virtual de hoje é um brinquedo de criança em comparação com a de uma civilização milhares de anos mais avançada que a nossa.

Isso significa que as mentes mais brilhantes deles podem ter decidido experimentar vidas imaginárias em vários mundos virtuais. Ele admite que esse pensamento é desalentador, mas não deixa de ser uma possibilidade. Pode até nos ajudar a aperfeiçoarmos nossa realidade virtual.

O QUE ELES QUEREM?

No filme *Matrix*, as máquinas assumem o comando e colocam os humanos em casulos, onde nos exploram como baterias para se energizarem. Só por isso elas nos mantêm vivos. Mas como uma única usina elétrica gera mais potência que os corpos de milhões de humanos, qualquer alienígena em busca de uma fonte de energia perceberá rapidamente que não precisa de baterias humanas (isso deve ter passado despercebido às máquinas soberanas de *Matrix*, mas esperemos que os alienígenas levem em consideração).

Outra possibilidade é que queiram nos comer. Isso foi explorado num episódio de *Além da imaginação*, em que alienígenas aterrissam na Terra prometendo benefícios tecnológicos e convidam voluntários para conhecer seu belo planeta. Mas os alienígenas esquecem aqui um livro chamado *Como servir homens*, que

os cientistas tentam decifrar a fim de descobrir as maravilhas que eles têm para compartilhar conosco, e descobrem ser um livro de culinária (mas como somos feitos de DNA e proteínas totalmente diferentes das deles, podem ter dificuldade para nos digerir).

Outra possibilidade é que queiram saquear todos os recursos e minerais valiosos da Terra. Esse argumento pode ter certo fundamento, mas, se eles são tão avançados ao ponto de viajar tranquilamente pelas estrelas, devem encontrar muitos planetas desabitados, com recursos que podem saquear sem se aborrecer com nativos indóceis. Do ponto de vista deles, seria uma perda de tempo tentar colonizar um planeta habitado quando há alternativas melhores.

Então, se os alienígenas não querem nos saquear nem escravizar, que perigo eles oferecem? Pense num veado na floresta. De quem ele deve ter mais medo: de um caçador cruel armado de espingarda, ou de um construtor tranqüilo com uma planta topográfica na mão? O caçador pode assustar o veado, mas só alguns veados são ameaçados por ele. Mais perigoso é o amável construtor porque o veado nem aparece no seu radar. Ele está tão concentrado em transformar a floresta num empreendimento imobiliário, que nem pensa no veado. Em vista disso, como seria de fato uma invasão?

Nos filmes de Hollywood, há uma falha gritante: os alienígenas estão sempre por volta de um século à nossa frente, e geralmente inventamos uma arma secreta ou exploramos uma fraqueza qualquer na estrutura deles para expulsá-los, como em *A invasão dos discos voadores*. Mas o diretor do Seti, dr. Seth Shostak, me disse que uma batalha com uma civilização alienígena avançada seria como uma luta entre o Bambi e o Godzilla.

Na realidade, os alienígenas devem ser milhares ou milhões de anos mais avançados em armamentos, então, de um modo geral, temos poucas chances de defesa. Mas talvez possamos aprender com os bárbaros que derrotaram o maior império militar da época, o Império Romano.

Os romanos eram mestres em engenharia, capazes de criar armas que arrasavam vilas, e de construir estradas para levar suprimentos a postos avançados do vasto império. Os bárbaros, recém-saídos de uma existência nômade, tinham poucas chances contra o Exército Romano e seu poder destruidor.

Mas a história registra que, à medida que se expandia, o império ia enfraquecendo, por lutar em muitas frentes, afundar-se em tratados e dispor de uma economia insuficiente para sustentar tudo, principalmente em vista do declínio gradual da população. Além disso, com escassez de homens, o império precisava recrutar jovens soldados bárbaros e promovê-los a posições de comando. Inevitavelmente, a tecnologia superior romana foi passando para os bárbaros, que, ao fim de algum tempo, começaram a dominar a própria tecnologia militar que os dominara.

Já no fim, o império enfraquecido por intrigas internas, falta de comida, guerras civis e um exército espalhado demais, deparou-se com bárbaros capazes de enfrentar o Exército Romano até deixá-lo inativo. Os saques de Roma entre 410 e 455 d.C. abriram o caminho para a queda final do império, em 476.

Da mesma forma, é provável que num primeiro momento os terráqueos não representem um perigo real no caso de uma invasão alienígena, mas com o tempo descobrirão os pontos fracos do exército invasor, suas fontes de energia, seus centros de comando, e principalmente de seus armamentos. A fim de controlar a população humana, os alienígenas terão que recrutar colaboradores. Isso resultará numa difusão da tecnologia entre os humanos.

Assim, um exército desordenado de humanos pode conseguir organizar um contra-ataque. Na estratégia militar oriental, como os ensinamentos clássicos de Sun Tzu em *A arte da guerra*, há um meio de derrotar um exército superior. Primeiro, permite-se que ele invada o território. Quando estiverem em terreno desconhecido e os soldados dispersos, é possível contra-atacar no ponto mais fraco deles.

Outra técnica é usar a força do inimigo contra ele. No judô, a principal estratégia é usar a seu favor o ímpeto do adversário. Você deixa o adversário atacar e nesse momento passa uma rasteira ou o derruba de repente, explorando a massa e energia dele. Quanto mais pesado, com mais força ele bate no chão. Talvez o único meio de lutar contra um exército superior seja deixar que invada o território, e então tentar descobrir seus segredos militares e conhecer seus armamentos, e depois usar as próprias armas e segredos do invasor contra ele.

Portanto, um exército alienígena não pode ser derrotado numa luta direta, mas terá que ir embora se não puder vencer e o custo de um empate for muito alto. O sucesso depende de privar o inimigo de uma vitória.

Acredito, porém, que é muito provável que os alienígenas sejam benevolentes e praticamente nos ignorem. Não temos nada a oferecer a eles. Se vierem, será mais por curiosidade, ou para fazer um reconhecimento. Como a curiosidade foi uma característica essencial para nos tornarmos inteligentes, é provável que qualquer espécie alienígena seja curiosa, e talvez queira nos analisar, mas não necessariamente fazer contato.

ENCONTRO COM UM ASTRONAUTA ALIENÍGENA

Ao contrário do que acontece nos filmes, provavelmente não teremos um encontro pessoal com criaturas alienígenas. Seria muito perigoso e desnecessário. Assim como enviamos *rovers* para explorar Marte, eles devem enviar substitutos orgânicos/mecânicos ou avatares, que podem suportar melhor as adversidades da

viagem interestelar. Dessa maneira, os “alienígenas” que encontrarmos nos gramados da Casa Branca podem não parecer em nada com seus senhores no planeta nativo. Os senhores projetarão sua consciência no espaço por meio de representantes.

Muito provavelmente, eles enviarão sondas robóticas à nossa lua, que é geologicamente estável, sem erosão. Essas sondas são autorreplicantes, isto é, criam uma fábrica e produzem, digamos, mil cópias delas (são chamadas sondas de Von Neumann, em homenagem ao matemático John von Neumann, que lançou as bases do computador digital; ele foi o primeiro matemático a considerar seriamente a questão das máquinas que se reproduzem). Essa segunda geração de sondas é enviada para outros sistemas estelares, onde cada uma delas cria uma terceira geração de outras mil sondas, totalizando um milhão. Depois elas se dispersam e criam outras fábricas, chegando ao total de um bilhão de sondas. Tendo início com apenas uma sonda, chegam a mil, um milhão, um bilhão. Em cinco gerações, um quatrilhão de sondas. Logo haverá uma esfera gigantesca se expandindo quase à velocidade da luz, contendo trilhões e trilhões de sondas, colonizando a galáxia em poucas centenas de milhares de anos.

Davies leva tão a sério essa ideia das sondas autorreplicantes de Von Neumann que já pediu financiamento para vasculhar a superfície da Lua à procura de vestígios de presença alienígena. Ele quer fazer uma varredura na Lua para detectar emissões de rádio ou anomalias de radiações que indiquem a presença de alienígenas, talvez milhões de anos atrás. Ele publicou um artigo com o dr. Robert Wagner no periódico científico *Acta Astronautica* pedindo que examinassem atentamente as fotos do Lunar Reconnaissance Orbiter com uma resolução de até 45 centímetros.

No artigo, diziam: “Embora haja uma probabilidade mínima de que uma tecnologia alienígena tenha deixado vestígios na Lua, na forma de artefatos ou modificação superficiais das características lunares, esse local tem a virtude de estar próximo”, e traços de uma tecnologia alienígena devem permanecer por longos períodos de tempo. Dado que não há erosão na Lua, marcas deixadas por alienígenas ainda seriam visíveis (da mesma forma que as pegadas deixadas por nossos astronautas nos anos 1970 podem, em princípio, permanecer lá por bilhões de anos).

Um problema é que a sonda de Von Neumann pode ser muito pequena. Nanossondas usam máquinas moleculares e MEMs, portanto, podem ser do tamanho de um pão de forma, ou menor, ele me disse (de fato, se uma sonda dessas pousasse num quintal aqui na Terra, o dono da casa talvez nem notasse).

Entretanto, esse método representa o meio mais eficaz de colonização da galáxia, usando o crescimento exponencial das sondas de Von Neumann. (É o mesmo meio pelo qual um vírus infecta nosso corpo. Começa com alguns vírus entrando nas células, em seguida sequestram a maquinaria reprodutiva das

células e as transformam em fábricas que criam mais vírus. Em duas semanas, um único vírus pode infectar trilhões de células, e começamos a espirrar.)

Se isso estiver correto, significa que nossa lua é o lugar mais propício para uma visita alienígena. E é também a base do filme *2001: Uma odisseia no espaço*, que até hoje apresenta o encontro mais plausível com uma civilização extraterrestre. No filme, uma sonda é colocada em nossa lua milhões de anos atrás, com o objetivo principal de observar a evolução da vida na Terra. Às vezes a sonda interfere e dá um impulso em nossa evolução. Essa informação é enviada para Júpiter, que é uma estação de retransmissão, e de lá segue para o planeta nativo de uma civilização alienígena muito antiga.

Do ponto de vista dessa civilização avançada, que pode analisar bilhões de sistemas estelares simultaneamente, vemos que eles podem escolher o sistema planetário que quiserem para colonizar. Dada a imensidão da galáxia, eles podem coletar dados e escolher à vontade os planetas e luas que prometem os melhores recursos. Da perspectiva deles, a Terra não deve ser muito atraente.

Os impérios do futuro serão impérios da mente.

– WINSTON CHURCHILL

Se continuarmos a desenvolver a tecnologia sem sabedoria ou prudência, o servo pode passar a ser o carrasco.

– GENERAL OMAR BRADLEY

15 OBSERVAÇÕES FINAIS

Em 2000, uma controvérsia feroz agitou a comunidade científica. Bill Joy, um dos fundadores da Sun Computers, escreveu um artigo inflamado denunciando a ameaça mortal representada pela tecnologia avançada. Num artigo publicado na revista *Wired* com o provocativo título “Por que o futuro não precisa de nós”, ele escreveu: “As tecnologias mais avançadas do século XXI – robótica, engenharia genética e nanotecnologia – ameaçam tornar os humanos uma espécie em extinção.” Esse artigo incendiário questionava a própria moralidade de centenas de cientistas trabalhando arduamente em seus laboratórios para o avanço da ciência. Ele pôs em questão a própria essência das pesquisas, afirmando que os benefícios da tecnologia eram ofuscados pela enorme ameaça que representava para a humanidade.

Joy descreveu uma distopia macabra, em que todas as nossas tecnologias conspiravam para destruir a civilização. E advertiu que três de nossas principais criações se voltarão contra nós:

- Um dia, os germes produzidos pela bioengenharia vão escapar do laboratório e devastar o mundo. Como não é possível recapturar essas formas de vida, elas podem se proliferar sem controle e desencadear uma epidemia mortal em todo o planeta, pior que as pestes da Idade Média. A biotecnologia pode até alterar a evolução humana, criando “várias espécies separadas e desiguais (...) que ameaçarão a noção de igualdade, que é a verdadeira pedra fundamental da democracia”.
- Um dia, os nanorrobôs podem enlouquecer e expelir grandes quantidades de “gosma cinzenta”, que asfixiariam todas as formas de vida. Como esses nanorrobôs “digerem” matéria comum para criar novas formas de matéria, os defeituosos podem se descontrolar e digerir boa parte da Terra. Joy escreveu: “Certamente, a gosma cinzenta seria um fim deprimente para a aventura humana na Terra, muito pior que o fogo ou o gelo, e pode se originar de um simples acidente de laboratório. Oops!”
- Um dia, os robôs vão dominar e substituir a humanidade. Ficarão tão inteligentes que vão simplesmente descartar os humanos. Seremos apenas uma nota de rodapé na história da civilização. “Os robôs nunca serão nossos filhos. (...) Nesse caminho, a humanidade pode estar perdida”, ele escreveu.

Joy afirmou que os perigos deflagrados por essas três tecnologias reduzirão a migalhas os perigos trazidos pela bomba atômica nos anos 1940, quando Einstein avisou que a tecnologia nuclear tinha o poder de destruir a civilização: “Tornou-se terrivelmente óbvio que nossa tecnologia superou nossa humanidade.” Mas a bomba atômica foi construída por um grande programa governamental, que podia ser estritamente regulamentado, ao passo que essas novas tecnologias são desenvolvidas por empresas privadas com uma regulamentação fraca, se é que têm alguma, argumenta Joy.

Ele concorda que, em curto prazo, essas tecnologias podem aliviar alguns sofrimentos. Mas, em longo prazo, os benefícios serão suplantados pelo fato de poderem desencadear um Armagedon científico capaz de aniquilar a raça humana.

Joy chegou ao ponto de acusar os cientistas de egoísmo e ingenuidade por tentarem criar uma sociedade melhor. Ele escreveu: “Uma boa sociedade e uma vida boa são parte da utopia tradicional. Uma vida boa envolve outras pessoas. Essa utopia tecnológica apela para ‘não vou ter doenças, não vou morrer, vou enxergar melhor e ser mais inteligente’, e por aí vai. Se disséssemos isso a Sócrates ou a Platão, eles dariam uma gargalhada.”

Ele conclui afirmando: “Penso que não é exagero dizer que estamos beirando a total perfeição do mal extremo, um mal cujas possibilidades vão muito além do que as armas de destruição em massa legaram aos Estados-nação.”

Qual é a conclusão? “Algo como a extinção”, ele adverte.

Como era de se esperar, esse artigo detonou controvérsias explosivas.

O artigo foi escrito há mais de uma década. Em termos de alta tecnologia, é uma vida inteira. Hoje, em retrospecto, é possível considerar certas previsões dele. Relendo o artigo numa perspectiva atualizada, vemos claramente que Bill Joy exagerou muitas das ameaças dessas tecnologias, mas também incitou os cientistas a encarar as consequências éticas, morais e sociais de seus trabalhos, o que é sempre bom.

E esse artigo deu origem a uma discussão sobre quem somos nós. Ao desvelar os segredos moleculares, genéticos e neurais do cérebro, não desumanizamos de certa forma a humanidade, reduzindo-a a um monte de átomos e neurônios? Se fizermos um mapa completo de todos os neurônios do cérebro e de todos os percursos neurais, não estaremos dissolvendo o mistério, a magia de quem somos nós?

UMA RESPOSTA A BILL JOY

Em retrospecto, as ameaças da robótica e da nanotecnologia estão mais distantes

do que Bill Joy supôs, e eu diria que, com as devidas advertências, podemos tomar uma série de contramedidas, como vetar pesquisas que levem à fabricação de robôs incontroláveis, inserir chips que os façam parar quando se tornarem perigosos, e criar dispositivos de emergência que imobilizem todos eles em caso de necessidade.

Mais imediata é a ameaça da biotecnologia, em que há um perigo real de vazamento de biogermes dos laboratórios. Na verdade, Ray Kurzweil e Bill Joy escreveram um artigo criticando a publicação do genoma completo do vírus da gripe espanhola de 1918, um dos germes mais letais da história moderna, que matou mais do que a Primeira Guerra Mundial. Os cientistas conseguiram remontar o vírus, já exterminado, examinando os corpos e sangue das vítimas e sequenciando seus genes. Depois, publicaram na internet.

Já existem formas de proteção contra o escape de vírus perigosos, mas ainda faltam medidas para fortalecer tais salvaguardas e aumentar a segurança. Por exemplo: se um novo vírus irromper em algum lugar distante, aqui na Terra, os cientistas precisam encontrar uma maneira rápida de isolar esse vírus num local remoto, sequenciar seus genes e preparar uma vacina para evitar a disseminação.

IMPLICAÇÕES PARA O FUTURO DA MENTE

Esse debate tem também um impacto direto no futuro da mente. No momento, a neurociência ainda é muito primitiva. Os cientistas podem ler e gravar em vídeo pensamentos simples do cérebro vivo, registrar algumas lembranças, conectar o cérebro a braços mecânicos, habilitar pacientes paralisados a controlar máquinas ao seu redor, silenciar regiões cerebrais específicas via magnetismo, e identificar nas doenças mentais as regiões do cérebro que funcionam mal.

Nas próximas décadas, porém, o poder da neurociência pode se tornar explosivo. Pesquisas em andamento estão no limiar de novas descobertas científicas que nos deixarão boquiabertos. Um dia será rotina controlar objetos usando apenas o poder da mente, gravar lembranças, curar doenças mentais, aumentar a inteligência, entender cada neurônio, criar cópias de dados do cérebro, e se comunicar por telepatia. O mundo do futuro será o mundo da mente.

Bill Joy não duvidava do potencial da tecnologia para aliviar o sofrimento humano. O que o horrorizou foi a possibilidade de indivíduos privilegiados levarem a uma cisão da humanidade. Em seu artigo, ele pinta uma distopia deprimente, em que uma pequena elite terá maior inteligência e aprimoramento dos processos mentais, enquanto as massas viverão na ignorância e na pobreza.

Ele temia que a raça humana se dividisse em duas, ou até que deixasse totalmente de ser humana.

Mas, como observamos, quase todas as tecnologias são no início muito caras e, portanto, exclusivamente para os ricos. Devido à produção em massa, à queda do custo dos computadores, à competição e ao transporte mais barato, as tecnologias acabam chegando aos pobres também. Essa foi a trajetória do fonógrafo, rádio, TV, PC, laptop e telefone celular.

Longe de criar um mundo de uns que têm e outros que não têm, a ciência é o motor da prosperidade. De todos os recursos que a humanidade utiliza, desde o início dos tempos, o mais forte e produtivo tem sido a ciência. A incrível riqueza que vemos por toda parte resulta diretamente da ciência. Para entendermos como a tecnologia reduz – e não acentua – os abismos sociais, basta lembrar como nossos ancestrais viviam até 1900. A expectativa de vida nos Estados Unidos era de 49 anos. Muitas crianças morriam na primeira infância. A comunicação com os vizinhos era feita por gritos na janela. O correio chegava a cavalo, quando chegava. A medicina era à base de poções mágicas. Os únicos tratamentos que funcionavam eram as amputações (sem anestesia) e a aplicação de morfina para aliviar a dor. A comida se estragava em poucos dias. Não havia saneamento. As doenças eram uma ameaça constante. E a economia só dava conta de alguns ricos e uma pequena classe média.

A tecnologia mudou tudo. Não precisamos mais caçar para comer; basta ir ao supermercado. Não precisamos carregar sacos pesados de mantimentos; basta levar até o carro. (Na verdade, a maior ameaça da tecnologia, que mata milhões de pessoas, não são robôs assassinos, nem nanorrobôs enlouquecidos – é nosso estilo de vida displicente, que criou níveis quase epidêmicos de diabetes, obesidade, doença cardíaca, câncer etc. E esse problema é autoinfligido.)

Vemos isso também em nível global. Nas últimas décadas, pela primeira vez na história, o mundo testemunhou centenas de milhões de pessoas saindo de uma pobreza opressiva. No panorama geral, vemos uma parcela significativa da raça humana deixando para trás as condições precárias de fazendas de subsistência para fazer parte da classe média.

As nações ocidentais levaram séculos para se industrializar e, no entanto, a China e a Índia fizeram isso em poucas décadas, graças à difusão da alta tecnologia. Contando com a tecnologia sem fio e a internet, essas nações podem ultrapassar países mais desenvolvidos que tiveram todo o trabalho de instalar redes elétricas nas cidades. Enquanto o Ocidente se debate com uma infraestrutura urbana antiga, decadente, as nações em desenvolvimento estão construindo cidades inteiras com uma brilhante tecnologia de última geração.

Quando eu era aluno de doutorado, meus colegas na China e na Índia tinham que esperar meses para receber publicações científicas pelo correio. Quase não tinham contato direto com cientistas e engenheiros ocidentais, porque poucos

podiam se dar ao luxo de viajar para cá. Isso travava muito o fluxo da tecnologia, que se movia a passos extremamente lentos naquelas nações. Hoje os cientistas podem ler os artigos de colegas assim que são publicados na internet, e colaborar com cientistas do mundo todo via e-mail. Isso acelerou muito o fluxo de informações. E com essa tecnologia, vêm o progresso e a prosperidade.

Ademais, não está claro por que um certo aprimoramento da inteligência provocará uma cisão catastrófica da raça humana, ainda que muitos não possam arcar com o custo dos procedimentos. Para a maioria, ser capaz de resolver equações matemáticas complexas ou ter memória perfeita não é garantia de um salário melhor, de maior respeito por parte dos colegas ou de maior popularidade entre o sexo oposto, que são incentivos para muita gente. O Princípio do Homem das Cavernas triunfa sobre cérebros turbinados.

Como diz o dr. Michael Gazzaniga: “A ideia de que estamos remexendo em nossos órgãos preocupa muita gente. E, afinal, o que faremos com a inteligência expandida? Vamos usá-la para solucionar problemas, ou só para fazer listas maiores de cartões de Natal...?”

Mas, como vimos no capítulo 5, trabalhadores desempregados podem se beneficiar, reduzindo drasticamente o tempo exigido para aumentar sua capacidade e dominar novas tecnologias. Isso pode não somente reduzir os problemas associados ao desemprego, mas também ter um impacto na economia mundial, tornando-a mais eficiente e adaptável a mudanças.

SABEDORIA E DEBATE DEMOCRÁTICO

Em resposta ao artigo de Joy, alguns críticos observaram que não se trata de uma luta entre os cientistas e a natureza, como diz o autor. Trata-se de um debate entre três partes: os cientistas, a natureza e a sociedade.

Os cientistas de computação John Brown e Paul Duguid responderam ao artigo dizendo: “As tecnologias – como a pólvora, a imprensa, a ferrovia, o telégrafo e a internet – causam mudanças profundas na sociedade. Por outro lado, os sistemas sociais – na forma de governos, sistemas de justiça, organizações formais e informais, movimentos sociais, redes de profissionais, comunidades locais, instituições financeiras, e tantos outros – modelam, moderam e redirecionam o poder bruto das tecnologias.”

Basta analisá-las do ponto de vista da sociedade e, por fim, cabe a nós adotar uma nova visão do futuro que incorpore todas as boas ideias.

Para mim, a fonte principal de sabedoria a esse respeito está num vigoroso debate democrático. Nas próximas décadas, o público será chamado a votar em várias questões científicas cruciais. A tecnologia não pode ser debatida no vácuo.

QUESTÕES FILOSÓFICAS

Enfim, alguns críticos afirmam que a ciência foi longe demais na revelação dos segredos da mente, uma revelação que desumaniza e degrada. Por que apreciar a descoberta de algo novo, aprender uma nova habilidade, divertir-se numa viagem de férias, se tudo pode ser reduzido a alguns neurotransmissores ativando alguns circuitos neurais?

Em outras palavras, assim como a astronomia nos reduziu a grãos insignificantes de poeira cósmica num universo indiferente, a neurociência nos reduziu a sinais elétricos circulando em circuitos neurais. Mas, isso é mesmo verdade?

Começamos nossa discussão destacando os dois maiores mistérios de toda a ciência: a mente e o universo. Ambos não têm apenas uma mesma história e narrativa, também compartilham de uma filosofia similar, e talvez do mesmo destino. A ciência, com todo o poder de investigar o centro de buracos negros e aterrissar em planetas distantes, deu origem a duas filosofias abrangentes sobre a mente e o universo: o princípio de Copérnico e o princípio antrópico. Ambos são compatíveis com tudo o que se conhece da ciência, mas são diametralmente opostos.

A primeira grande filosofia, o princípio de Copérnico, nasceu com a descoberta do telescópio, mais de 400 anos atrás, e afirma que a humanidade não tem uma posição privilegiada no universo. Essa ideia, aparentemente simples, derrubou mitos milenares e filosofias arraigadas.

Desde o conto bíblico de Adão e Eva exilados do Jardim do Éden por terem provado da maçã do conhecimento, houve uma série de derrotas humilhantes. Primeiro, o telescópio de Galileu mostrou que a Terra não era o centro do sistema solar, e sim o Sol. Essa ideia foi sobrepujada quando se soube que o sistema solar era apenas um cisco na Via Láctea, em órbita a cerca de 30 mil anos-luz do centro da galáxia. Nos anos 1920, Edwin Hubble descobriu inúmeras galáxias. De repente, o universo ficou bilhões de vezes maior. Agora o telescópio espacial Hubble revela a presença de até 100 bilhões de galáxias no universo visível. Nossa Via Láctea foi reduzida a um pontinho numa arena cósmica muitíssimo maior.

Teorias cosmológicas mais recentes rebaixam cada vez mais a posição da humanidade no universo. A teoria do universo inflacionário diz que nosso universo visível, com seus 100 bilhões de galáxias, é do tamanho de um furinho de agulha num universo muito maior, tão inflado que a maior parte da luz de regiões mais distantes ainda não teve tempo de chegar até nós. Na vastidão do espaço, há muitos lugares que não conseguimos avistar com telescópios, e aonde nunca chegaremos porque não podemos viajar mais rápido que a luz. E se a teoria das cordas (minha especialidade) estiver correta, significa que o universo inteiro

coexiste com outros universos num hiperespaço de onze dimensões. Portanto, o espaço tridimensional não é o ponto final. A verdadeira arena dos fenômenos físicos é o multiverso de universos, cheio de universos flutuando como bolhas de sabão.

O autor de ficção científica Douglas Adams tentou resumir o sentimento de estar sempre sendo destronado, inventando em *O guia do mochileiro das galáxias* o Vórtex da Perspectiva Total, criado para enlouquecer as pessoas. Quando alguém entra numa câmara, vê um mapa gigantesco de todo o universo, com uma setinha minúscula, quase invisível, apontando “Você está aqui”.

Portanto, o princípio de Copérnico indica que somos um entulho cósmico insignificante vagando sem rumo entre as estrelas. Por outro lado, todos os últimos dados cosmológicos são compatíveis com uma teoria que defende a filosofia oposta: o princípio antrópico.

Essa teoria afirma que o universo é compatível com a vida. Mais uma vez, essa afirmação aparentemente simples tem implicações profundas. Por um lado, é impossível negar que existe vida no universo, mas é claro que as forças do universo devem ser calibradas com precisão para tornar a vida possível. Como disse o físico Freeman Dyson: “O universo parecia saber que estávamos vindo.”

Por exemplo: se a força nuclear fosse só um pouquinho maior, o Sol teria explodido bilhões de anos atrás, antes que o DNA pudesse brotar. Se a força nuclear fosse só um pouquinho menor, o Sol jamais poderia ter entrado em ignição, e também não estaríamos aqui.

Da mesma forma, se a gravidade fosse mais forte, o universo teria se contraído até o Grande Colapso bilhões de anos atrás, e teríamos morrido torrados. Se fosse um pouquinho mais fraca, o universo teria se expandido tão depressa, atingindo o Grande Congelamento, e todos teríamos morrido de frio.

Esse ajuste se estende a todos os átomos do corpo. A física diz que somos feitos de poeira de estrelas, que os átomos que vemos em tudo à nossa volta foram forjados no calor de uma estrela. Somos literalmente filhos das estrelas.

Mas as reações nucleares que queimaram hidrogênio para criar os elementos mais pesados no nosso corpo são muito complexas, e poderiam ter dado errado por diversos motivos. Nesse caso, teria sido impossível criar os elementos mais pesados do nosso corpo, e os átomos de DNA e da vida não existiriam.

Em outras palavras, a vida é preciosa e é um milagre.

Há tantos parâmetros que precisaram de ajustes que alguns afirmam que não pode ser coincidência. A forma fraca do princípio antrópico sugere que a existência da vida obriga os parâmetros físicos do universo a serem definidos com a maior precisão. A forma forte do princípio antrópico vai mais além, afirmando que Deus, ou algum projetista, precisou criar um universo “bem certinho” para tornar a vida possível.

O debate entre o princípio de Copérnico e o princípio antrópico reverbera na neurociência. Por exemplo: alguns afirmam que os humanos podem ser reduzidos a átomos, moléculas e neurônios, consequentemente, não há um lugar diferenciado para a humanidade no universo.

David Eagleman diz: “Esse *você* que seus amigos conhecem e amam só existe enquanto os transistores e parafusos do seu cérebro estiverem no lugar. Se você não acredita, vá à enfermaria da ala neurológica de um hospital. Uma lesão, até numa pequena parte do cérebro, pode levar a uma perda terrível de capacidades específicas: a capacidade de nomear animais, ouvir música, evitar comportamentos de risco, distinguir cores, tomar decisões simples.”

Ao que parece, o cérebro não funciona sem todos os “transistores e parafusos”. Ele conclui, dizendo: “Nossa realidade depende das condições de nossa biologia.”

Então, por um lado, nosso lugar no universo parece diminuir se podemos ser reduzidos, como robôs, a porcas e parafusos (biológicos). Somos apenas um *wetware* ^[6] rodando um software chamado mente, nada mais, nada menos. Nossos pensamentos, desejos, esperanças e aspirações podem ser reduzidos a impulsos elétricos circulando numa região do córtex pré-frontal. Isso é o princípio de Copérnico aplicado à mente.

Mas o princípio antrópico também pode ser aplicado à mente, e chegamos à conclusão oposta. Diz simplesmente que as condições do universo tornam possível a consciência, ainda que seja extraordinariamente difícil criar uma mente a partir de eventos aleatórios. O grande biólogo vitoriano Thomas Huxley disse: “Algo tão notável como um estado de consciência ter surgido em decorrência de uma irritação do tecido nervoso é tão inexplicável quanto a aparição do Gênio quando Aladim esfregou a lâmpada.”

Além disso, muitos astrônomos acreditam que, embora um dia possamos encontrar vida em outros planetas, será provavelmente a vida microbiana que reinou em nossos mares por bilhões de anos. Em vez de grandes cidades e impérios, vamos achar somente oceanos de microrganismos à deriva.

Quando entrevistei o biólogo de Harvard Stephen Jay Gould, hoje já falecido, sobre este assunto, ele explicou seu pensamento da seguinte maneira: Se fôssemos criar um planeta gêmeo da Terra como ela era há 4,5 bilhões de anos, ele se tornaria idêntico ao que a Terra é hoje depois que se passassem 4,5 bilhões de anos? Provavelmente, não. Há uma grande probabilidade de que o DNA e a vida jamais brotassem, e uma probabilidade ainda maior de que a vida inteligente dotada de consciência jamais emergisse dos pântanos.

Gould escreveu: “O *Homo sapiens* é um galinho [da árvore da vida]. (...) Mas

esse galinho, para o bem e para o mal, desenvolveu a mais extraordinária qualidade de toda a história da vida multicelular desde a explosão cambriana (há 500 milhões de anos). Inventamos a consciência com todas as suas consequências, de Hamlet a Hiroshima.”

De fato, no curso da história da Terra, foram muitas as vezes em que a vida inteligente quase se extinguiu. Além das extinções em massa que aniquilaram os dinossauros e quase todas as formas de vida, os humanos enfrentaram outras ameaças de extinção. Por exemplo: os humanos são geneticamente relacionados entre si em um grau considerável, muito mais próximo que dois animais da mesma espécie. Embora cada ser humano tenha uma aparência diferente, nossos genes e nossa química interna contam uma história distinta. Na verdade, quaisquer dois humanos são tão próximos em termos de genes que podemos calcular matematicamente quando uma “Eva genética” e um “Adão genético” deram origem a toda a raça humana. E podemos até calcular quantos de nós existíamos no passado.

Os números são impressionantes. A genética mostra que havia somente centenas ou milhares de humanos vivos entre 70 e 100 mil anos atrás, e que eles deram à luz toda a raça humana (uma teoria sustenta que a explosão do vulcão Toba, na Indonésia, há cerca de 70 mil anos, provocou uma queda de temperatura tão drástica que eliminou a maior parte da raça humana, deixando uns poucos para povoar a Terra). Desse pequeno grupo de humanos vieram os aventureiros e exploradores que começaram a colonizar o planeta.

Como já foi dito, em muitos momentos da história da Terra a vida inteligente esteve perto de chegar ao fim. É um milagre termos sobrevivido. Podemos também concluir que, embora possa existir vida em outros planetas, só pode haver vida consciente numa parcela mínima deles. Temos que valorizar muito a consciência encontrada na Terra. É a mais alta forma de complexidade conhecida no universo, e provavelmente a mais rara.

Às vezes, contemplando o porvir da raça humana, me vejo às voltas com a forte possibilidade de sua autodestruição. Apesar de erupções vulcânicas e terremotos terem o poder de condenar a raça humana à morte, o medo maior é de calamidades produzidas pelo homem, como guerras nucleares ou germes criados pela bioengenharia. Se assim for, talvez a única forma de vida consciente na Via Láctea será extinta. Creio que será uma tragédia não só para nós, mas também para o universo. Temos absoluta certeza de que somos conscientes, mas não compreendemos a longa e tortuosa sequência de eventos biológicos que conseguiram tornar isso possível. O psicólogo Steven Pinker disse: “Defendo que nada dá mais sentido à vida do que compreender que cada momento de consciência é um dom precioso e frágil.”

O MILAGRE DA CONSCIÊNCIA

Por fim, uma crítica à ciência diz que entender algo significa remover seu mistério e magia. Ao levantar o véu que esconde os segredos da mente, a ciência a torna mais corriqueira e banal. Entretanto, quanto mais aprendo sobre a complexidade do cérebro, mais fico deslumbrado com o fato de nosso pescoço carregar o objeto mais sofisticado que conhecemos no universo. Como diz David Eagleman: “Que assombrosa obra de arte é o cérebro, e como somos sortudos de pertencer a uma geração que tem a tecnologia e a vontade para chamar a atenção para essa obra de arte. É a coisa mais maravilhosa que descobrimos no universo, e somos nós.” Em vez de minimizar o deslumbramento, aprender sobre o cérebro só o aumenta.

Há mais de dois mil anos Sócrates já dizia: “Conhecer a ti mesmo é o começo da sabedoria.” Temos uma longa jornada até chegar lá.

6. *Wetware* refere-se à ideia de um “computador biológico”, uma vez que a água é elemento comum a todos os seres vivos (da junção de *wet*, “molhado”, com *ware*, em alusão ao sufixo das palavras *software* e *hardware*). (N. do R.T.)

APÊNDICE

CONSCIÊNCIA QUÂNTICA?

Apesar de todos os avanços miraculosos em varreduras cerebrais e alta tecnologia, alguns afirmam que nunca compreenderemos o segredo da consciência, que está além da nossa tecnologia. De fato, na visão delas a consciência é mais do que átomos, moléculas e neurônios, e determina a natureza da própria realidade. Para elas, a consciência é a entidade fundamental, a origem da criação do mundo material. E como prova, referem-se a um dos maiores paradoxos de toda a ciência, que desafia nossa definição de realidade: o paradoxo do Gato de Schrödinger. Até hoje sem consenso universal, cientistas premiados assumem posições divergentes sobre a questão. O que está em jogo é nada menos que a natureza da realidade e do pensamento.

O paradoxo do Gato de Schrödinger chega aos fundamentos da mecânica quântica, o campo que torna possíveis os lasers, as varreduras por IRM, o rádio e a TV, a eletrônica moderna, o GPS e as telecomunicações, dos quais depende a economia mundial. Muitas previsões da teoria quântica passaram em testes com margem de precisão de uma parte em cem bilhões.

Passei toda minha carreira profissional trabalhando com a teoria quântica. E ainda acho que é como areia movediça. É uma sensação desconfortável saber que o trabalho de minha vida inteira se baseia numa teoria que tem os próprios fundamentos baseados em um paradoxo.

Esse debate foi iniciado pelo físico austríaco Erwin Schrödinger, um dos pais da teoria quântica. Ele estava tentando explicar o estranho comportamento de elétrons, que pareciam exibir propriedades tanto de partículas quanto de ondas. Como pode um elétron, uma partícula pontual, ter dois comportamentos divergentes? Às vezes os elétrons agiam como uma partícula, criando rastros bem definidos numa câmara de nuvem. Outras vezes, os elétrons agiam como uma onda, passando por furinhos e criando padrões de interferência de onda, como as na superfície de um lago.

Em 1925, Schrödinger apresentou sua célebre equação de onda, que leva seu nome e é uma das equações mais importantes jamais escritas. Foi um sucesso imediato e lhe valeu o Prêmio Nobel em 1933. A equação de Schrödinger descreveu precisamente o comportamento do elétron como onda e, quando aplicada ao átomo de hidrogênio, explicou suas estranhas propriedades. Milagrosamente, a equação podia ser aplicada também a qualquer átomo, explicando quase todas as características dos elementos da tabela periódica. Parecia que toda a química (e, portanto, toda a biologia) era apenas soluções

dessa equação de onda. Alguns físicos chegaram a afirmar que todo o universo, incluindo todas as estrelas, planetas e até nós, eram apenas soluções dessa equação.

Mas então os físicos começaram a se fazer uma pergunta problemática que ressoa até hoje: Se o elétron é descrito por uma equação de onda, então o que é a ondulação?

Em 1927, Werner Heisenberg propôs um novo princípio que dividiu os físicos. O célebre princípio da incerteza de Heisenberg afirma que não se podem saber ao mesmo tempo, com certeza, a localização e o momento (o produto de sua massa pela sua velocidade) do elétron. Essa incerteza não derivava da imprecisão dos instrumentos, mas era inerente à própria física. Nem Deus, nem ser celestial algum poderia saber com precisão a localização e o momento de um elétron.

Assim, a equação de onda de Schrödinger na verdade descrevia a probabilidade de se encontrar o elétron. Os cientistas tinham passado milhares de anos tentando a duras penas eliminar de seus trabalhos o acaso e as probabilidades, e agora Heisenberg os trazia de volta.

A nova filosofia pode ser resumida no seguinte: o elétron é uma partícula pontual, mas a probabilidade de encontrá-lo é dada por uma onda. E essa onda obedece à equação de Schrödinger, dando origem ao princípio da incerteza.

A comunidade da física partiu-se ao meio. De um lado, físicos como Niels Bohr, Werner Heisenberg e a maioria dos físicos atômicos logo adotaram a nova formulação. Quase todos os dias eles anunciavam descobertas quanto à compreensão das propriedades da matéria. Prêmios Nobel eram dados a físicos quânticos como Oscars. A mecânica quântica estava se tornando um livro de receitas. Não era preciso ser um mestre em física para fornecer contribuições estelares — bastava seguir as receitas dadas pela mecânica quântica para fazer descobertas espantosas.

Do outro lado, físicos já idosos, premiados com o Nobel, como Albert Einstein, Erwin Schrödinger e Louis de Broglie levantaram objeções filosóficas. Schrödinger, pai do trabalho desencadeador de todo esse processo, resmungou que, se soubesse que sua equação acabaria introduzindo a probabilidade dentro da física, jamais a teria criado.

Os físicos travaram um debate de oitenta anos, que persiste até hoje. Se Einstein afirmava que “Deus não joga dados com o mundo”, ouviu-se Niels Bohr responder: “Pare de dizer a Deus o que fazer.”

Em 1935, para arruinar os físicos quânticos de uma vez por todas, Schrödinger propôs seu famoso problema do gato. Coloque o gato numa caixa selada, com um frasco de gás venenoso. Dentro da caixa, há um pouco de urânio. O átomo de urânio é instável e emite partículas que podem ser detectadas por um contador Geiger. O contador aciona um martelo, que cai e quebra o vidro, liberando o gás,

que pode matar o gato.

Como se descreve o gato? Um físico quântico diria que o átomo de urânio é descrito por uma onda, que tanto pode decair como não decair. Portanto, é preciso adicionar essas duas ondas juntas. Se o urânio decair, o gato morre, e isso é descrito por uma onda. Se o urânio não decair, então o gato vive, o que também é descrito por uma onda. Para descrever o gato, então, é preciso somar a onda do gato morto com a onda do gato vivo.

Isso significa que o gato não está morto nem vivo! O gato está num mundo do nunca, entre a vida e a morte, na soma da onda que descreve um gato morto com a onda do gato vivo.

Esse é o xis do problema, que vem reverberando pelo mundo da física por quase um século. E como se resolve esse paradoxo? Há pelo menos três maneiras (e centenas de variações dessas três).

A primeira é a interpretação original de Copenhague, proposta por Bohr e Heisenberg, citada nos livros didáticos do mundo inteiro (e por onde começo quando ensino mecânica quântica). Nessa interpretação, para se determinar o estado do gato é preciso abrir a caixa e fazer uma medição. A onda do gato (que é a soma do gato morto com o gato vivo) “colapsa” em uma onda só, de modo que agora se sabe que o gato está vivo (ou morto). Assim, a observação determina a existência e o estado do gato. O processo de medição é, portanto, responsável pela mágica dissolução de duas ondas em uma.

Einstein odiava isso. Durante séculos os cientistas combateram o chamado “solipsismo” ou “idealismo subjetivo”, ou seja, os objetos só existem se alguém os observa. Apenas a mente é real – o mundo material só existe como ideias na mente. Portanto, dizem os solipsistas (como o bispo George Berkeley), se uma árvore cai na floresta sem ninguém por perto observando, então talvez a árvore nem tenha caído. Einstein, que achava tudo isso pura bobagem, propôs uma teoria oposta, chamada “realidade objetiva”, dizendo simplesmente que o universo existe em um estado único, definido, independentemente de qualquer observação humana. É o conceito comum da maioria das pessoas.

A realidade objetiva remonta a Isaac Newton. Nesse cenário, o átomo e as partículas subatômicas são como bolinhas de aço, que existem em pontos definidos no espaço e tempo. Não há ambiguidade nem acaso para localizar a posição dessas bolinhas, que seguem as leis do movimento. A realidade objetiva teve grande sucesso na descrição dos movimentos dos planetas, estrelas e galáxias. Usando a relatividade, essa ideia também descreve buracos negros e o universo em expansão. Mas há um lugar em que ela infelizmente falha: dentro do átomo.

Os físicos clássicos, como Newton e Einstein, pensavam que a realidade objetiva finalmente eliminava o solipsismo da física. O jornalista Walter Lippmann resumiu: “A novidade radical da ciência moderna está precisamente

na rejeição da crença (...) de que as forças que movem as estrelas e os átomos dependem das preferências do coração humano.”

Mas a mecânica quântica deixou entrar mais uma vez na física uma nova forma de solipsismo. Nesse quadro, antes que seja observada, uma árvore pode existir em qualquer estado possível (por exemplo, brotando, queimada, transformada em serragem, palitos de dentes, caída). Mas quando se olha para ela, a onda subitamente colapsa, e se vê a forma de árvore. Os solipsistas anteriores falavam em árvores caídas ou não caídas. Os novos solipsistas quânticos introduziram *todos* os estados possíveis da árvore.

Isso foi demais para Einstein. Em sua casa, perguntava aos convidados: “A Lua existe porque um ratinho olha para ela?” Para um físico quântico, em certo sentido a resposta pode ser sim.

Einstein e seus colegas desafiaram Bohr, perguntando como o micromundo quântico (com gatos simultaneamente vivos e mortos) coexistia com o mundo do senso comum que todos vemos? A resposta era a existência de uma “parede” separando nosso mundo do mundo atômico. De um lado da parede, rege o senso comum. Do outro, rege a teoria quântica. Mesmo que a parede se mova, os resultados continuam os mesmos.

Essa interpretação, por estranha que seja, é ensinada há oitenta anos pelos físicos quânticos. Recentemente, recaíram algumas dúvidas sobre a interpretação de Copenhague. Hoje temos a nanotecnologia, que permite manipular átomos um a um, à vontade. Na tela de um microscópio de tunelamento com varredura, os átomos parecem bolinhas de tênis difusas. (Para a BBC, tive a oportunidade de visitar o Almaden Lab da IBM em San Jose, na Califórnia, e, com uma minúscula sonda, realmente empurrei átomos, um de cada vez. Hoje é possível brincar com os átomos, antes considerados tão pequenos que jamais seriam enxergados.)

Como já discutimos, a Idade do Silício está chegando lentamente ao fim, e há quem diga que os transistores de tal material serão substituídos pelos moleculares. Se assim for, no futuro, os paradoxos da teoria quântica poderão estar dentro de cada computador. A economia mundial pode vir a depender desses paradoxos.

CONSCIÊNCIA CÓSMICA E UNIVERSOS MÚLTIPLOS

Há duas interpretações alternativas para o paradoxo do gato que nos levam aos domínios mais estranhos da ciência: o reino de Deus e os múltiplos universos.

Em 1967, a segunda resolução do problema do gato foi formulada pelo ganhador do Prêmio Nobel Eugene Wigner, cujo trabalho foi crucial para os fundamentos da mecânica quântica e também para a construção da bomba

atômica. Ele disse que só uma pessoa consciente pode fazer uma observação que colapse a função de onda. Mas quem disse que essa pessoa existe? Não se pode separar o observador do observado, então talvez esta pessoa também esteja morta ou viva. Portanto, deve haver uma nova função de onda que inclua o gato e o observador. Para garantir que o observador esteja vivo, é necessário um segundo observador para vigiar o primeiro. Esse segundo observador, chamado de “amigo de Wigner”, é necessário para vigiar o primeiro, de modo que todas as ondas colapsem. Mas como se sabe que o segundo observador está vivo? O segundo observador deve ser incluído em uma função de onda ainda mais extensa, para garantir que ele está vivo, mas isso pode continuar indefinidamente. Como é necessário um número infinito de “amigos” para colapsar a função de onda anterior garantindo que estão vivos, acaba sendo necessária alguma forma de “consciência cósmica”, ou Deus.

Wigner concluiu que “não era possível formular as leis (ou a teoria quântica) de modo coerente sem fazer referência à consciência”. No fim da vida, ele chegou a se interessar pela filosofia vedanta do hinduísmo.

Nessa abordagem, Deus ou alguma consciência eterna observa a todos, colapsando nossas funções de onda, de modo que podemos dizer que estamos vivos. Essa interpretação fornece os mesmos resultados físicos que a de Copenhague; portanto, essa teoria não pode ser refutada. Mas sugere que a consciência é a entidade fundamental no universo, mais fundamental que os átomos. O mundo material pode ir e vir, mas a consciência permanece como o elemento definidor, significando que a consciência, de certo modo, cria a realidade. A própria existência dos átomos que nos rodeiam se baseia em nossa capacidade de vê-los e tocá-los.

Neste ponto, é importante notar que, para algumas pessoas, se a consciência determina a existência, então a consciência pode controlar a existência, talvez por meditação. Elas pensam que podemos criar a realidade de acordo com nossos desejos. Esse pensamento, por mais atraente que possa soar, contraria a mecânica quântica. Na física quântica, a consciência faz observações e assim determina o estado da realidade, mas a consciência não pode escolher com antecedência qual estado de realidade de fato existe. A mecânica quântica permite determinar a chance de encontrar um estado, mas não podemos curvar a realidade aos nossos desejos. Por exemplo, num jogo de pôquer, é possível calcular matematicamente as chances de se obter um *royal straight flush*. Contudo, isso não significa que, de alguma forma, seja possível controlar as cartas para recebê-lo. Não é possível escolher universos, assim como não podemos controlar se o gato está morto ou vivo.

A terceira forma de resolver o paradoxo do gato é a interpretação de Everett, ou de muitos mundos, proposta em 1957 por Hugh Everett. É a teoria mais estranha de todas, afirmando que o universo está constantemente se dividindo em um multiverso de universos. Em um universo, temos um gato morto. Em outro, o gato está vivo. Essa abordagem pode ser resumida dizendo que as funções de onda nunca colapsam, apenas se dividem. A teoria de muitos mundos de Everett difere da interpretação de Copenhague somente por eliminar a última suposição – o colapso da função de onda. De certo modo, é a formulação mais simples da mecânica quântica, mas também a mais perturbadora.

Há profundas consequências dessa terceira abordagem. Significa que todos os universos possíveis talvez existam, mesmo os bizarros, aparentemente impossíveis (contudo, quanto mais bizarro o universo, mais improvável).

Assim, as pessoas que morreram no nosso universo ainda estão vivas em outro. E essas pessoas mortas insistem que o universo delas é o correto, enquanto o nosso (onde estão mortas) é falso. Mas se esses “fantasmas” de pessoas mortas ainda estão vivos em algum lugar, por que não conseguimos encontrá-los? Por que não podemos tocar nesses mundos paralelos? (Por incrível que pareça, dentro dessa perspectiva, Elvis ainda está vivo em um desses universos.)

Além disso, alguns dos universos podem estar mortos, sem vida nenhuma, enquanto outros podem ser exatamente como o nosso, exceto por uma diferença crucial. Por exemplo, a colisão de um único raio cósmico é um evento quântico minúsculo. Mas o que acontece se esse raio cósmico cai na mãe de Adolf Hitler, e o menino morre por aborto? Então, um evento quântico minúsculo, a colisão de um único raio cósmico, causa a divisão do universo ao meio. Em um universo, a Segunda Guerra Mundial nunca existiu e sessenta milhões de pessoas não morreram. No outro, temos a devastação da Segunda Guerra Mundial. Esses dois universos vão se afastando até ficarem bem diferentes, embora tenham sido separados por um evento quântico minúsculo.

Esse fenômeno foi explorado pelo escritor de ficção científica Philip K. Dick no romance *O homem do castelo alto*, em que um universo paralelo se abre por causa de um único evento: uma bala é disparada e Franklin Roosevelt morre assassinado. Esse acontecimento faz com que os Estados Unidos não se preparem para a Segunda Guerra Mundial. Os nazistas e japoneses vencem e acabam dividindo os Estados Unidos ao meio.

Porém, a bala dispara ou falha dependendo de uma faísca microscópica atingir ou não a pólvora que, por sua vez, depende de reações moleculares complexas envolvendo movimentos de elétrons. Assim, talvez as flutuações quânticas na pólvora determinem se a arma dispara ou não, o que determina, por sua vez, se são os aliados ou os nazistas que saem vitoriosos da Segunda Guerra Mundial.

Assim, não há “parede” separando o mundo macro do quântico. As

características bizarras da teoria quântica podem penetrar em nosso mundo do “senso comum”. Essas funções de onda nunca colapsam – continuam se dividindo eternamente em realidades paralelas. A criação de universos alternativos não para nunca. Os paradoxos do micromundo (isto é, estar vivo e morto simultaneamente, estar em dois lugares ao mesmo tempo, desaparecer e reaparecer em outro local) agora entram também em nosso mundo.

Mas se a função de onda está sempre se dividindo, criando universos inteiramente novos nesse processo, então por que não podemos visitá-los?

Steven Weinberg, ganhador do prêmio Nobel, faz uma comparação com ouvir rádio na sala de casa. Há centenas de ondas de rádio simultâneas enchendo a sala, vindas do mundo inteiro, mas o rádio está sintonizado em uma só frequência. Em outras palavras, o rádio está em “decoerência” com todas as outras estações (coerência significa que todas as ondas estão vibrando perfeitamente em uníssono, como num feixe de raios laser; a decoerência se dá quando essas ondas ficam fora de fase, portanto, não vibram mais em uníssono). Todas as outras frequências continuam existindo, mas o rádio não as capta porque elas não estão vibrando na mesma frequência que foi sintonizada. Elas se desacoplaram, isto é, estão em decoerência com a frequência selecionada.

Do mesmo modo, as funções de onda do gato morto e do gato vivo entram em decoerência com o passar do tempo. As implicações são bastante perturbadoras. Na sala de casa, podemos coexistir com ondas de dinossauros, piratas, alienígenas do espaço, monstros. Mas ficamos tranquilos, sem ter a mínima noção de estarmos compartilhando o mesmo lugar com esses habitantes do espaço quântico, porque nossos átomos não estão mais vibrando em uníssono com eles. Esses universos paralelos não existem numa terra do nunca distante. Eles existem na sala da nossa casa.

A entrada em um desses mundos paralelos é chamada de “salto quântico” e é o artifício preferido da ficção científica. Para entrar em um universo paralelo, é preciso pular para dentro dele. (Houve até uma série de TV chamada *Sliders – Dimensões paralelas*, em que os personagens entravam e saíam de universos paralelos. A série começa com um menino lendo um livro. Na verdade, é um livro meu, *Hiperespaço*, mas não assumo nenhuma responsabilidade pela física por trás da série.)

E não é tão simples pular de um universo para outro. Um dos problemas que às vezes passamos aos alunos de doutorado é calcular a probabilidade de saltar contra uma parede de tijolos e ir parar do outro lado dela. O resultado é um balde de água fria. Seria preciso esperar mais que o tempo de vida do universo para vivenciar um salto através de uma parede de tijolos.

Quando me olho no espelho, não me vejo como realmente sou. Em primeiro lugar, me vejo a cerca de um bilionésimo de segundo atrás, que é o tempo necessário para um feixe de luz sair do meu rosto, atingir o espelho e entrar nos meus olhos. Segundo, a imagem que vejo é na verdade uma média de bilhões e mais bilhões de funções de onda. Essa média certamente se parece com minha imagem, mas não é exata. Ao redor, múltiplas imagens minhas ondulam em todas as direções. Estou constantemente cercado de universos alternativos, ramificando-se eternamente em diferentes mundos, mas a possibilidade de deslizar de um para outro é tão ínfima que a mecânica newtoniana parece estar correta.

Neste ponto, alguns perguntam: Por que os cientistas não fazem logo um experimento para determinar qual interpretação é válida? Se fizermos um experimento com um elétron, todas as três interpretações darão o mesmo resultado. Portanto, as três são interpretações sérias e viáveis da mecânica quântica, de acordo com a mesma teoria quântica subjacente. A diferença está na forma de explicar os resultados.

Daqui a centenas de anos, físicos e filósofos talvez ainda estejam debatendo essa questão, sem chegar a nenhuma conclusão, porque todas as três interpretações dão os mesmos resultados físicos. Mas talvez haja um caminho nesse debate filosófico que passe pelo cérebro, e trata-se do livre-arbítrio que, por sua vez, afeta os alicerces morais da sociedade humana.

LIVRE-ARBÍTRIO

Toda a nossa civilização se baseia no conceito de livre-arbítrio, com impactos nos conceitos de recompensa, punição e responsabilidade pessoal. Mas existe realmente o livre-arbítrio? Ou é uma forma engenhosa de manter a sociedade unida, mesmo violando princípios científicos? A controvérsia chega ao cerne da própria mecânica quântica.

É seguro dizer que cada vez mais cientistas chegam à conclusão de que o livre-arbítrio não existe, ao menos no sentido usual. Se certos comportamentos estranhos podem ser atribuídos a defeitos identificados no cérebro, então uma pessoa não é cientificamente responsável pelos crimes que possa cometer. Ela pode ser perigosa demais para andar à solta nas ruas, e deve ser trancafiada em algum tipo de instituição; mas punir alguém por ter um derrame ou tumor no cérebro é injusto, dizem os cientistas. A pessoa precisa é de ajuda médica e psicológica. Talvez o dano cerebral possa ser tratado (por exemplo, removendo um tumor) e ela possa se tornar um membro produtivo da sociedade.

Quando entrevistei o dr. Simon Baron-Cohen, psicólogo da Universidade de

Cambridge, ele me disse que muitos (mas não todos) assassinos patológicos têm alguma anomalia cerebral. A varredura do cérebro mostra que eles não têm empatia pelo sofrimento dos outros e de fato podem até sentir prazer em vê-los sofrer (nesses indivíduos, a amígdala e o núcleo accumbens, o centro do prazer, se iluminam quando assistem a vídeos de pessoas sentindo dor).

A conclusão que se pode tirar daí é que essas pessoas não são realmente responsáveis por seus atos hediondos mas, mesmo assim, deveriam ser isoladas da sociedade. Elas precisam de ajuda, não punição, devido ao problema que têm no cérebro. Em certo sentido, não agem com livre-arbítrio quando cometem um crime.

Um experimento conduzido pelo dr. Benjamin Libet, em 1985, lança dúvidas sobre a própria existência de livre-arbítrio. Digamos que se peça às cobaias que observem um relógio e prestem atenção na posição exata do ponteiro quando decidirem mover um dedo. Usando varreduras por EEG, pode-se detectar exatamente quando o cérebro toma essa decisão. Comparando os dois instantes, encontra-se um descompasso. As varreduras do EEG mostram que o cérebro, na verdade, tomou a decisão cerca de 300 milissegundos antes que a pessoa fique ciente.

Isso significa que, de certo modo, o livre-arbítrio é falso. As decisões são tomadas com antecedência pelo cérebro, sem intervenção consciente, e depois o cérebro tenta encobrir o feito (como é seu costume), afirmando que a decisão foi consciente. O dr. Michael Sweeney conclui que “os resultados de Libet indicam que o cérebro sabe o que a pessoa vai decidir *antes* que ela o faça. (...) O mundo precisa reavaliar não só a ideia da divisão dos movimentos em voluntários e involuntários, mas também a própria noção de livre-arbítrio”.

Tudo isso parece indicar que o livre-arbítrio, a base da sociedade, é uma ficção, uma ilusão criada pelo cérebro esquerdo. Então somos senhores do nosso destino, ou meros peões numa fraude perpetrada pelo cérebro?

Há várias maneiras de abordar essa antiga questão. O livre-arbítrio contraria uma filosofia chamada determinismo, que simplesmente considera todos os eventos como sendo determinados por leis da física. De acordo com o próprio Newton, o universo é uma espécie de relógio, funcionando desde o início dos tempos, de acordo com as leis do movimento. Assim, todos os eventos são previsíveis.

A questão é: fazemos parte desse relógio? Todas as nossas ações são também determinadas? Essas questões têm implicações filosóficas e teológicas. Por exemplo, quase todas as religiões aderem a alguma forma de determinismo e predestinação. Como Deus é onipotente, onisciente e onipresente, Ele conhece o futuro e o determina com antecedência. Mesmo antes de você nascer, Ele sabe se você vai para o Céu ou para o Inferno.

A Igreja Católica se dividiu exatamente por essa questão, durante a revolução

protestante. Segundo a doutrina católica na época, era possível alterar o destino obtendo uma indulgência, geralmente em troca de generosas doações financeiras à Igreja. Em outras palavras, o determinismo podia ser alterado de acordo com o tamanho da carteira. Martinho Lutero destacou especificamente a corrupção da Igreja com as indulgências em suas 95 Teses, que afixou na porta de uma igreja, em 1517, deflagrando a Reforma Protestante. Essa foi uma das razões básicas para a cisão da Igreja, que causou a morte de milhões e espalhou destruição em regiões inteiras da Europa.

Porém, depois de 1925, a incerteza foi introduzida na física pela mecânica quântica. De repente, tudo se tornou incerto; só se podia calcular a probabilidade. Nesse sentido, talvez o livre-arbítrio realmente exista, e é uma manifestação da mecânica quântica. Por isso alguns afirmam que a teoria quântica restabelece o conceito de livre-arbítrio. Os deterministas contra-atacam, declarando que os efeitos quânticos são extremamente pequenos (no nível dos átomos), pequenos demais para dar conta do livre-arbítrio dos seres humanos.

A situação hoje, na verdade, está bastante confusa. Talvez a pergunta “Existe livre-arbítrio?” seja parecida com “O que é vida?”. A descoberta do DNA tornou obsoleta essa pergunta sobre a vida. Agora sabemos que a questão tem muitas camadas e complexidades. Talvez o mesmo se aplique ao livre-arbítrio, e há muitos tipos.

Se assim for, a própria definição de “livre-arbítrio” se torna ambígua. Por exemplo, uma forma de definir o livre-arbítrio é verificando se o comportamento pode ser previsto. Se existe o livre-arbítrio, então o comportamento não pode ser determinado com antecedência. Digamos que você está assistindo a um filme. A trama é completamente determinada, sem nenhum livre-arbítrio. Então, o filme é completamente previsível. Mas nosso mundo não pode ser como um filme, por duas razões. A primeira é a teoria quântica, como vimos. O filme representa apenas uma linha de tempo possível. A segunda é a teoria do caos. Embora a física clássica diga que todos os movimentos dos átomos são completamente determinados e previsíveis, na prática é impossível prever tais movimentos, porque há muitos átomos envolvidos. A mais leve perturbação em um só átomo pode ter um efeito cascata nos movimentos, criando enormes perturbações.

Pense na previsão do tempo, por exemplo. Em princípio, sabendo o comportamento de cada átomo no ar, seria possível prever o tempo de daqui a cem anos, caso tenhamos um computador grande o suficiente. Mas na prática, isso é impossível. Depois de apenas algumas horas, o tempo fica tão turbulento e complexo que qualquer simulação no computador torna-se inútil.

Isso cria o chamado “efeito borboleta” – até um bater de asas de uma borboleta é capaz de provocar pequenas ondulações na atmosfera, que crescem e podem causar uma tempestade. Assim, se até o bater de asas de uma borboleta

pode causar tempestades, acreditar que se consegue prever o tempo com precisão é uma fantasia.

Voltemos ao experimento que Stephen Jay Gould me descreveu, sobre o pensamento. Ele me pediu para imaginar a Terra há 4,5 bilhões de anos, quando nasceu. Agora imagine que você conseguiu criar uma cópia idêntica da Terra e a deixou evoluir. Será que os humanos também teriam surgido nessa outra Terra depois de 4,5 bilhões de anos?

É natural pensar que, pelos efeitos quânticos ou pela natureza caótica do tempo e dos oceanos, a evolução jamais levaria a humanidade às mesmas características das criaturas daquela versão da Terra. Enfim, parece que a combinação de incerteza e caos torna impossível um mundo perfeitamente determinístico.

O CÉREBRO QUÂNTICO

Esse debate também afeta a engenharia reversa do cérebro. Se for possível aplicá-la com êxito, tornando o cérebro um conjunto de transistores, então o cérebro é determinístico e previsível. Para qualquer pergunta, ele sempre repetirá a mesma resposta. Dessa forma, os computadores são determinísticos, porque sempre dão a mesma resposta para todas as perguntas.

Parece que temos um problema. Por um lado, a mecânica quântica e a teoria do caos afirmam que o universo não é previsível e, portanto, é preciso existir o livre-arbítrio. Mas um cérebro construído por engenharia reversa, feito de transistores, por definição será previsível. Como esse cérebro teoricamente é idêntico ao cérebro vivo, então o cérebro humano também é determinístico e não há livre-arbítrio. É claro que isso contradiz a primeira afirmação.

Uma minoria de cientistas declara que não se consegue fazer uma engenharia reversa autêntica do cérebro, nem mesmo criar uma máquina que realmente pense, por causa da teoria quântica. Argumentam que o cérebro é um dispositivo quântico, não apenas um conjunto de transistores. Assim, esse projeto está condenado ao fracasso. Nesse campo atua o físico de Oxford, autoridade na teoria da relatividade de Einstein, dr. Roger Penrose. Ele afirma que os processos quânticos podem ser os responsáveis pela consciência do cérebro humano. Penrose começou dizendo que o matemático Kurt Gödel provou que a aritmética é incompleta; isto é, existem afirmações verdadeiras na aritmética que não podem ser provadas usando os axiomas da própria disciplina.

Analogamente, não só a matemática é incompleta, mas também a física. Ele conclui dizendo que o cérebro é basicamente um dispositivo mecânico quântico e que há problemas que nenhuma máquina resolve, devido ao teorema da

incompletude de Gödel. Mas os humanos conseguem entender esses enigmas usando a intuição.

Do mesmo modo, o cérebro criado por engenharia reversa, por mais complexo que seja, não passa de um conjunto de transistores e fios. Em um sistema determinista como esse, pode-se prever exatamente seu comportamento futuro porque as leis do movimento são bem conhecidas. Em um sistema quântico, porém, o sistema em si é inerentemente imprevisível. Só é possível calcular a probabilidade de ocorrer um evento por causa do princípio da incerteza.

Se chegarem à conclusão de que o cérebro da engenharia reversa não consegue reproduzir o comportamento humano, então os cientistas podem ser obrigados a admitir que há forças imprevisíveis em ação (isto é, efeitos quânticos dentro do cérebro). Penrose argumenta que dentro do neurônio há minúsculas estruturas, chamadas microtúbulos, onde reinam os processos quânticos.

No momento, não há consenso sobre esse problema. A julgar pelas reações às ideias de Penrose logo que foram apresentadas, é prudente dizer que a maior parte da comunidade científica vê sua abordagem com ceticismo. Mas a ciência não é conduzida por pesquisas de popularidade, e seus avanços dependem de teorias que podem ser testadas, reproduzidas, ou provadas como sendo falsas.

De minha parte, acredito que os transistores não conseguem modelar de fato todos os comportamentos dos neurônios, que executam cálculos tanto analógicos como digitais. Sabemos que os neurônios são confusos. Podem vazar, errar, envelhecer, morrer, e são sensíveis ao ambiente. Para mim, isso indica que um conjunto de transistores consegue apenas se aproximar do comportamento dos neurônios. Por exemplo, vimos na discussão sobre a física do cérebro que se o axônio do neurônio ficar mais fino começa a vazar, e passa a não realizar as reações químicas como deveria. Alguns desses vazamentos e erros nos disparos são devidos a efeitos quânticos. À medida que imaginamos os neurônios mais finos, mais densos e mais rápidos, tais efeitos ficam mais óbvios. Isso significa que mesmo os neurônios normais estão sujeitos a vazamentos e instabilidades, e esses problemas existem tanto em relação à física clássica quanto à mecânica quântica.

Concluindo, um robô construído por engenharia reversa será uma aproximação boa, mas não perfeita, do cérebro humano. Discordando de Penrose, acho que é possível criar um robô determinístico feito de transistores, que aparente ter consciência, mas sem nenhum livre-arbítrio. E vai passar no teste de Turing. Mas acho que haverá diferenças entre um robô desse tipo e um ser humano, por causa desses minúsculos efeitos quânticos.

Por fim, acho que o livre-arbítrio provavelmente existe, mas não aquele defendido pelos individualistas radicais, que afirmam ser senhores absolutos de seu destino. O cérebro é influenciado por milhares de fatores inconscientes que

nos predispõem a fazer certas escolhas com antecedência, mesmo acreditando tê-las feito por vontade própria. Isso não significa necessariamente que somos atores em um filme que pode ser rebobinado à vontade. O final do filme ainda não foi escrito, e o determinismo radical é destruído por uma combinação sutil de efeitos quânticos com a teoria do caos. Afinal, ainda somos senhores do nosso destino.

INTRODUÇÃO

17 Seria preciso viajar: Para ter uma noção disso, é preciso definir “complexo” em termos da quantidade total de informações que podem ser armazenadas. A rival mais próxima do cérebro pode ser a informação contida no nosso DNA. Há três bilhões de pares de bases em nosso DNA, cada um contendo um dos quatro ácidos nucleicos, rotulados A, T, C, G. Portanto, a quantidade total de informações que podemos armazenar no DNA é quatro elevado à potência de três bilhões. Mas o cérebro pode armazenar muito mais informações em seus cem bilhões de neurônios, que podem disparar ou não. Consequentemente, o número de estados iniciais possíveis no cérebro humano é dois elevado à potência de cem bilhões. E enquanto o DNA é estático, os estados do cérebro mudam a cada poucos milissegundos. Um mero pensamento pode conter cem gerações de disparos neurais. Logo, em cem gerações, o número de pensamentos possíveis é dois elevado a cem bilhões, e tudo isso elevado à potência de cem. Mas nosso cérebro dispara continuamente, dia e noite, computando sem cessar. Portanto, o número total de pensamentos possíveis em n gerações é dois elevado à potência de cem bilhões, e tudo isso elevado à potência n , o que é um resultado astronômico. Logo, a quantidade de informações que pode ser armazenada no cérebro excede em muito a que é armazenada em nosso DNA. Na verdade, é a maior quantidade de informações que podemos armazenar em nosso sistema solar, e até, possivelmente, em nosso setor da Via Láctea.

21 “*Os insights mais valiosos*”: Boleyn-Fitzgerald, p. 89.

21 “*Todas essas questões que os filósofos*”: Boleyn-Fitzgerald, p. 137.

1. DESVENDANDO A MENTE

31 **Ficou semiconsciente durante semanas:** Ver Sweeney, pp. 207-8.

31 **O dr. John Harlow, o médico que tratou:** Carter, p. 24.

33 **No ano de 43 d.C. há registros:** Horstman, p. 87.

34 **“Parecia... que eu estava na porta:** Carter, p. 28.

48-49 **O CÉREBRO TRANSPARENTE:** *New York Times*, 10 de abril de 2013, p. 1.

53 **“as emoções não são sentimentos”:** Carter, p. 83.

53 **a mente é mais parecida com uma “sociedade de mentes”:** Entrevista com dr. Minsky para a série da BBC *Visions of the Future*, fevereiro de 2007. E também entrevista para o programa de rádio *Science Fantastic*, novembro de 2009.

53 **a consciência é como uma tempestade:** Entrevista com dr. Pinker em setembro de 2003, para o programa de rádio *Exploration*.

53 **“o sentimento intuitivo que temos de que existe um ‘eu’”:** Pinker, “The Riddle of Knowing You’re Here”, em *Your Brain: A User’s Guide* (Nova York: Time Inc. Specials, 2011).

53 **A consciência consiste em:** Boleyn-Fitzgerald, p. 111.

57 **“de fato um sistema consciente”:** Carter, p. 52.

57 **perguntei como podem ser feitos experimentos:** Entrevista com dr. Michael Gazzaniga em setembro de 2012 para o programa de rádio *Science Fantastic*.

58 **“As implicações possíveis”:** Carter, p. 53.

58 **“O que acontece quando essa pessoa morre?”:** Boleyn-Fitzgerald, p. 119.

58 **um jovem rei que herda o trono:** Entrevista com dr. David Eagleman em maio de 2012 para o programa de rádio *Science Fantastic*.

59 **“pessoas chamadas Denise ou Denis”:** Eagleman, p. 63.

59 **“pelo menos 15% das mulheres”:** Eagleman, p. 43.

2. CONSCIÊNCIA – O PONTO DE VISTA DE UM FÍSICO

60 “**Não conseguimos ver a luz ultravioleta**”: Pinker, *How the Mind Works*, pp. 561-65.

61 “**Todo mundo sabe o que é consciência**”: *Biological Bulletin* 215, n.3 (dezembro de 2008): 216.

64 **Explicaremos isso nas Notas**: O Nível II de consciência pode ser calculado contando o total de ciclos de feedback distintos nas interações de um animal com os membros de sua espécie. Numa aproximação grosseira, o Nível II de consciência de um dado animal pode ser estimado multiplicando-se a quantidade de outros animais no bando pela quantidade de emoções ou gestos distintos que o animal usa para se comunicar. Mas essa primeira estimativa requer considerações.

Por exemplo, animais como o lince são sociais, embora sejam caçadores solitários, o que pode dar a entender que só há um animal no bando. Mas isso só vale durante a caçada. Na época do acasalamento, os lincos cumprem rituais complexos, o que deve ser levado em conta para o Nível II de consciência.

Além disso, quando as fêmeas dão à luz, a ninhada precisa ser cuidada e alimentada, aumentando consequentemente o número de interações sociais. Assim, mesmo no caso de caçadores solitários, a quantidade de animais contados nas interações não é igual a um, e a quantidade total de ciclos distintos de feedback pode ser bem grande.

E se o número de lobos numa alcateia diminui, pode parecer que a classificação atribuída ao Nível II de consciência também diminui. No entanto, para considerar essas variações, precisamos introduzir o conceito de grau médio do Nível II para toda a espécie, bem como um Nível II específico de consciência para cada um dos animais.

O grau médio do Nível II de consciência da espécie não muda se o bando diminuir, porque vale para toda a espécie, mas o grau do Nível II para um indivíduo muda (pois mede consciência e atividade mental).

Aplicado aos humanos, o grau médio do Nível II deve levar em conta o número de Dunbar, igual a 150, que representa aproximadamente o número de pessoas em nosso grupo social com quem podemos manter contato. Assim, o grau do Nível II de consciência dos humanos, como uma espécie, seria a quantidade total de emoções e gestos distintos que usamos para nos comunicar multiplicada pelo número de Dunbar, 150. (Cada indivíduo pode ter diferentes graus no Nível II de consciência, já que seus círculos sociais e

modos de interagir com eles podem variar consideravelmente.)

É preciso considerar que certos organismos no Nível I de consciência (como insetos e répteis) podem manifestar comportamentos sociais. Quando esbarram umas nas outras, as formigas trocam informações por meio de odores de substâncias químicas, enquanto que as abelhas dançam para comunicar a localização de canteiros. Os répteis têm até um sistema límbico primitivo. Mas em geral não manifestam emoções.

65 **“A diferença entre o homem”**: Gazzaniga, p. 27.

65 **“A maior conquista do cérebro humano”**: Gilbert, p. 5.

65 **“área 10 (a camada granular interna IV)”**: Gazzaniga, p. 20.

67 **o macho fica confuso porque quer**: Eagleman, p. 144.

74 **“Prevejo que os neurônios espelho”**: Brockman, p. xiii.

77 **O biólogo dr. Carl Zimmer escreve**: Bloom, p. 51.

77 **“na maior parte do tempo, devaneamos”**: Bloom, p. 51.

78 **Perguntei a uma pessoa que poderia**: Entrevista com dr. Michael Gazzaniga em setembro de 2012 para o programa de rádio *Science Fantastic*.

78 **“é o hemisfério esquerdo”**: Gazzaniga, p. 85.

3. TELEPATIA – UM DOCE PELO QUE VOCÊ ESTÁ PENSANDO

- 84 **De fato, num recente “Next 5 in 5”:** <http://www.ibm.com/5in5>.
- 85 **Tive o prazer de conhecer o laboratório:** Entrevista com dr. Gallant em 11 de julho de 2012, na Universidade da Califórnia em Berkeley. Entrevista com dr. Gallant também para o programa de rádio *Science Fantastic*.
- 86 **“Isso é um salto muito importante”:** *Berkeleyan Newsletter*, 22 de setembro de 2011, <http://newscenter.berkeley.edu/2011/09/22/brain-movies>.
- 87 **“Tomando duzentos vóxeis”:** Brockman, p. 236.
- 88 **o dr. Brian Pasley e seus colegas:** Visita ao laboratório do dr. Pasley em 11 de julho de 2012, na Universidade da Califórnia, Berkeley.
- 89 **Resultados semelhantes foram obtidos:** The Brain Institute, Universidade de Utah, Salt Lake City, <http://brain.utah.edu>.
- 90 **Isso pode ter aplicação para artistas:** <http://io9.com/5423338/a-device-that-lets-you-type-with-your-mind>.
- 91 **Segundo seus representantes:** <http://news.discovery.com/tech/type-with-your-mind-110309.html>.
- 92 **explorado pelo dr. David Poeppel:** *Discover Magazine Presents the Brain*, primavera de 2012, p. 43.
- 93 **Em 1993, na Alemanha, o dr. Bernhard Blümich:** *Scientific American*, novembro de 2008, p. 68.
- 94 **A única justificativa para sua existência:** Garreau, pp. 23-24.
- 95 **Uma vez almocei com um ex-diretor:** Simpósio sobre o futuro da ciência, patrocinado pelo Science Fiction Channel no Chabot Space and Science Center, Oakland, Califórnia, em maio de 2004.
- 96 **Em outra ocasião:** Conferência em Anaheim, Califórnia, abril de 2009.
- 97 **“Imagine se soldados”:** Garreau, p. 22.
- 98 **“o que ele está fazendo é gastar”:** Ibid., p. 19.
- 99 **Quando perguntei ao dr. Nishimoto:** Visita ao laboratório do dr. Gallant, na Universidade da Califórnia em Berkeley, em 11 de julho de 2012.
- 100 **“Há questões éticas”:** <http://vitals.nbcnews.com/news/2012/01/31/10281528-words-from-brain-waves-may-let-scientists-read-your-mind>.

4. TELECINESIA – A MENTE CONTROLANDO A MATÉRIA

- 102 **“Eu adoraria ter”**: *New York Times*, 17 de maio de 2012, p. A17, e http://www.nbcnews.com/id/47447302/ns/health-health_care/t/paralyzed-woman-gets-robotic-arm.htm.
- 102 **“Pegamos um pequeno sensor”**: Entrevista com dr. John Donoghue em novembro de 2009 para o programa de rádio *Science Fantastic*.
- 103 **Somente nos Estados Unidos, mais de 200 mil**: Centers for Disease Control and Prevention, Washington, D.C. <http://www.cdc.gov/traumaticbraininjury/scifacts.html>.
- 104 **Quando o macaco queria mover o braço**: <http://www.northwestern.edu/newscenter/stories/2012/04/miller-paralyzed-technology.html>.
- 104 **“Estamos interceptando os sinais elétricos”**: <http://www.northwestern.edu/newscenter/stories/2012/04/miller-paralyzed-technology.html>.
- 105 **Mais de 1.300 combatentes retornaram**: <http://www.darpa.mil/program/revolutionizing-prosthetics>. *60 Minutes* da CBS, edição de 30 de dezembro de 2012. *60 Minutes* da CBS, edição de 30 de dezembro de 2012.
- 105 **“Acharam que estávamos loucos”**: Ibid.
- 105 **Jan foi ao programa 60 Minutes**: Ibid.
- 106 **“Haverá um ecossistema inteiro”**: *Wall Street Journal*, 29 de maio de 2012.
- 106 **Mas talvez a maior novidade**: Entrevista com dr. Nicolelis em abril de 2011 para o programa de rádio *Science Fantastic*.
- 107 **FUSÃO INTELIGENTE DE MÃOS E MENTES**: *New York Times*, 13 de março de 2013, http://huffingtonpost.com/2013/2/28/mind-melds-brain-communication_n_2781609.html. Ver também *Huffington Post*, 28 de fevereiro de 2013, <http://nytimes.com/2013/03/01/science/new-research-suggests-two-rat-brains-can-be-linked.html>. Ver também *Huffington Post*, 28 de fevereiro de 2013, .
- 108 **Em 2013, outro passo importante**: *USA Today*, 8 de agosto de 2013, p. 1D.
- 112 **Cerca de dez anos atrás**: Entrevista com dr. Nicolelis em abril de 2011.
- 112 **“para que não fique nada”**: Para uma discussão completa sobre o

exoesqueleto, ver Nicolelis, pp. 303-7.

- 116 **a Honda Corporation construiu:** <http://www.asimo.honda.com>. E também entrevista com os criadores do Asimo em abril de 2007 para a série de TV da BBC *Visions of the Future*. E também entrevista com os criadores do Asimo em abril de 2007 para a série de TV da BBC *Visions of the Future*.
- 117 **Depois, domina-se a técnica de:**
<http://discovermagazine.com/2007/may/review-test-driving-the-future>.
- 117 **Com o pensamento, o paciente:** *Discover*, 9 de dezembro de 2011,
<http://discovermagazine.com/2011/dec/09-mind-over-motor-controlling-robots-with-your-thoughts>.
- 117 **“Provavelmente seremos capazes de operar”:** Nicolelis, p. 315.
- 122 **Vi uma demonstração dessa tecnologia:** Entrevista com cientistas no Carnegie Mellon em agosto de 2010 para a série *Sci Fi Science* do canal de TV Discovery/Science.

5. MEMÓRIAS E PENSAMENTOS FEITOS SOB MEDIDA

- 126 “As peças só se encaixaram”: Wade, p. 89.
- 128 Até o momento, os cientistas identificaram: Ibid., p. 91.
- 128 Por exemplo: o dr. Antonio Damasio: Damasio, pp. 130-53.
- 129 Um fragmento de memória pode vibrar: Wade, p. 232.
- 129 “Se não for feito”: <http://newscientist.com/article/dn3488>.
- 130 “Liga o botão”: http://www.eurekalert.org/pub_releases/2011-06/uosc-rmr061211.php.
- 130 “implantes que aumentem a competência”:
<http://hplusmagazine.com/2009/03/18/artificial-hippocampus>.
- 130 Com tanta coisa em jogo, não é de surpreender que:
http://articles.washingtonpost.com/2013-07-12/national/40863765_1_brain-cells-mice-new-memories.
- 134 Se codificar a memória: Isso levanta a questão de se pombo-correio, pássaros migratórios, baleias etc. têm memória de longo prazo, dado que podem migrar centenas ou milhares de quilômetros em busca de comida e de locais para terem seus filhotes. A ciência sabe muito pouco sobre isso, mas acredita-se que a memória de longo prazo desses animais se baseia em localizar pontos de referência ao longo do caminho, e não na memória de um acontecimento passado. Ou seja, eles não utilizam a memória de algo que aconteceu para simular o futuro, sua memória de longo prazo consiste numa série de marcações. Parece que apenas os humanos são capazes de usar a memória de longo prazo para simular o futuro.
- 134 “o objetivo da memória”: Michael Lemonick, “Your Brain: A User’s Guide”, *Time*, dezembro de 2011, p. 78.
- 135 “Pode-se pensar nisso como”: http://www.scidai.ly/videos/2007/0710-brain_scans_of_the_future.ht.
- 135 seu estudo prova uma “resposta possível”:
<http://www.sciencedaily.com/videos/2007/0710>.
- 136 “A ideia é que o equipamento”: *New York Times*, 12 de setembro de 2012, p. A18.
- 136 “Provavelmente levaremos várias décadas”:

<http://www.tgdaily.com/general-sciences-features/58736-artificial-cerebellum-restores-rats-brain-function>.

- 137 Atualmente 5,3 milhões de norte-americanos: Alzheimer's Foundation of America, <http://www.alzfdn.org>.
- 139 “Isso reforça a noção”: ScienceDaily.com, outubro de 2009, <http://www.sciencedaily.com/releases/2009/10/091019122647.htm>.
- 139 “Nunca conseguiremos torná-lo um matemático”: Ibid.
- 140 “Isso indica que essas moscas”: Wade, p. 113.
- 140 Tal resultado não se restringe: Ibid.
- 140 “Podemos dar uma razão biológica”: Ibid., p. 114.
- 142 quanto mais as proteínas CREB: Bloom, p. 244.
- 144 “O propranolol se instala na célula nervosa”: SATI e-News, 28 de junho de 2007, <http://www.my sati.com/enews/june2007/ptsd.htm>.
- 144 Um relatório concluiu que: Boleyn-Fitzgerald, p. 104.
- 145 “Por mais que nossos rompimentos”: Ibid.
- 145 “deve privá-lo de morfina”: Ibid., p. 105.
- 145 “Se outros trabalhos confirmarem”: Ibid., p. 106.
- 149 “Cada uma dessas gravações perenes”: Nicolelis, p. 318.
- 149 “O esquecimento é o processo natural”: *New Scientist*, 12 de março de 2003, <http://www.newscientist.com/article/dn3488>.

6. O CÉREBRO DE EINSTEIN E A EXPANSÃO DE NOSSA INTELIGÊNCIA

152 “se deixado levar pelo momento”:

<http://abcnews.go.com/blogs/headlines/2012/03/einsteins-brain-arrives-in-london-after-odd-journey>.

154 “Sempre afirmei que”: Gould, p. 109.

155 “O cérebro humano continua tendo ‘plasticidade’”:

www.sciencedaily.com/releases/2011/12/111208257120.htm.

156 “O quadro extraído desses estudos”: Gladwell, p. 40.

157 **Cinco anos depois, Terman deu início:** Ver C. K. Holahan e R. R. Sears, *The Gifted Group in Later Maturity* (Stanford, CA : Stanford University Press, 1995).

158 “as notas na escola”: Boleyn-Fitzgerald, p. 48.

159 “os testes não medem motivação”: Sweeney, p. 26.

159 **Os pilotos com maior pontuação:** Bloom, p. 12.

159 “O hemisfério esquerdo é responsável”: Ibid., p. 15.

163 **O dr. Darold Treffert, de Wisconsin:** <http://www.daroldtreffert.com>.

163 **Em 45 segundos, ele respondeu:** Tammet, p. 4.

164 **Tive o prazer de entrevistar:** Entrevista com Daniel Tammet em outubro de 2007 para o programa de rádio *Science Fantastic*.

165 “Nosso estudo confirma”: *Science Daily*, março de 2012,

<http://www.sciencedaily.com/releases/2012/03/120322100313.htm>.

165 **O cérebro de Kim Peek** Matéria da Associated Press, 8 de novembro de 2004, <http://www.space.com>.

166 **Em 1998, dr. Bruce Miller:** *Neurology* 51 (outubro de 1998): pp. 978-82. Ver também http://www.wisconsinmedicalsociety.org/professional/savant-syndrome/resources/articles/the_acquired-savant.

167 **Além dos savants:** Sweeney, p. 252.

167 **essa ideia já foi testada:** Center of the Mind, Sydney, Austrália, <http://www.centreforthemind.com>.

168 **Em outro experimento, o dr. R. L. Young:** R. L. Young, M. C. Ridding e T. L. Morrell, “Switching Skills on by Turning Off Part of the Brain”,

- 168 “Quando aplicada aos lobos pré-frontais”: Sweeney, p. 311.
- 169 Até recentemente, pensava-se: *Science Daily*, maio de 2012, <http://www.sciencedaily.com/releases/2012/05/120509180113.htm>.
- 170 “Os savants têm uma alta capacidade”: Ibid.
- 172 Em 2007, houve um grande avanço: Sweeney, p. 294.
- 172 “As pesquisas de células-tronco”: Sweeney, p. 295.
- 172 Os cientistas se concentraram em alguns genes: Katherine S. Pollard, “What Makes Us Different”, *Scientific American Special Collectors Edition* (inverno de 2013): 31-35.
- 173 “Agarrei a oportunidade de participar”: Ibid.
- 174 “Com meu mentor David Haussler”: Ibid.
- 177 Um gene desses foi descoberto: *TG Daily*, 15 de novembro de 2012. <http://www.tgdaily.com/general-sciences-features/67503-new-found-gene-separates-man-from-apes>.
- 179 Muitas teorias, remontando a Charles Darwin: Ver, por exemplo, Gazzaniga, *Human: The Science Behind What Makes Us Unique*.
- 181 “Nas primeiras centenas de milhões de anos”: Gilbert, p. 15.
- 183 “Os neurônios da massa cinzenta cortical”: Douglas Fox, “The Limits of Intelligence”, *Scientific American*, julho de 2011, p. 43.
- 183 “mãe de todas as limitações”: Ibid., p. 42.

7. EM SEUS SONHOS

- 193 **E complementou o registro:** C. Hall e R. Van de Castle, *The Content Analysis of Dreams*, Nova York: Appleton-Century-Crofts, 1966.
- 196 **Quando o entrevistei, dr. Hobson:** Entrevista com dr. Allan Hobson em julho de 2012 para o programa de rádio *Science Fantastic*.
- 197 **Estudos mostram que é possível:** Wade, p. 229.
- 198 **Yukiyasu Kamitani, cientista chefe do ATR:** *New Scientist*, 12 de dezembro de 2008, <http://www.newscientist.com/article/dn16267-mindreading-software-could-record-your-dreams.html>.
- 198 **Quando visitei seu laboratório:** Visita ao laboratório do dr. Gallant em 11 de julho de 2012.
- 200 **“Nossos sonhos, portanto, não são”:** *Science Daily*, 28 de outubro de 2011, <http://www.sciencedaily.com/releases/2011/11/11028113626.htm>.
- 201 **Protótipos de lentes de contato:** Ver o trabalho do dr. Babak Parviz, <http://www.wearable-technologies.com/262>.

8. A MENTE PODE SER CONTROLADA?

203 **Um touro enfurecido adentra:** Miguel Nicolelis, *Beyond Boundaries*, Nova York Henry Holt, 2011, pp. 228-32.

206 **A histeria da Guerra Fria:** “Project MKULTRA, the CIA’s Program of Research into Behavioral Modification. Joint Hearings Before the Select Committee on Human Resources, U.S. Senate, 95th Congress, First Session”, Government Printing Office, 8 de agosto de 1977, Washington D.C., http://www.nytimes.com/packages/pdf/national/13inmate_ProjectM “CIA Diz Ter Encontrado Mais Documentos Secretos sobre Controle Comportamental”, *New York Times*, 3 de setembro de 1977; “Registros Governamentais de Controle da Mente da MKUltra e Bluebird/Artichoke”, <http://wanttoknow.info/mindcontrol.shtml>; “Comissão para Estudo de Operações Governamentais sobre Atividades da Inteligência Militar e no Exterior”. Relatório do Comitê dos Representantes de Igrejas n. 94-755, 94º Congresso, 2ª Sessão, p. 392, Government Printing Office, Washington, D.C., 1976; “Project MKULTRA, the CIA’s Program of Research in Behavioral Modification”, <http://scribd.com/doc/75512716/Project-MKUltra-The-CIA-s-Program-of-Research-in-Behavioral-Modification>.

207 “**grande potencial de desenvolvimento**”: Rose, p. 292.

208 “**impossibilidade neurocientífica**”: Ibid., p. 293.

209 “**É provavelmente significativo que**”: “Hypnosis and Intelligence,” Black Vault Freedom of Information Act Archive, 2008, <http://documents.theblackvault.com/documents/mindcontrol/hypnosisinin>

210 **Para avaliar a extensão desse problema:** Boleyn-Fitzgerald, p. 57.

211 **Drogas como o LSD:** Sweeney, p. 200.

212 “**É a primeira vez que vemos**”: Boleyn-Fitzgerald, p. 58.

213 “**Se quisermos desligar**”:
<http://www.nytimes.com/2011/05/17/science/17optics.html>.

214 “**Ao enviar informação dos sensores**”: *New York Times*, 17 de março de 2011, <http://nytimes.com/2011/05/17/science/17optics.html>.

9. ESTADOS ALTERADOS DE CONSCIÊNCIA

- 218 **“Uma parte dos profetas”**: Eagleman, p. 207.
- 219 **“Às vezes é um Deus próprio”**: Boleyn-Fitzgerald, p. 122.
- 219 **“Finalmente entendi tudo mesmo”**: Ramachandran, p. 280.
- 219 **“Durante os três minutos de estimulação”**: David Biello, *Scientific American*, p. 41, .
- 220 **Para testar essas ideias**: Ibid., p. 42.
- 221 **“Embora os ateus possam argumentar”**: Ibid., p. 45.
- 221 **“Se um ateu”**: Ibid., p. 44.
- 222 **Uma teoria sustenta que a doença de Parkinson**: Sweeney, p. 166.
- 223 **“Os neurônios conectados à sensação”**: Ibid., p. 90.
- 225 **“o cérebro vai continuar fazendo o que faz”**: Ibid., p. 165.
- 226 **“As varreduras cerebrais levaram pesquisadores”**: Ibid., p. 208.
- 226 **“Se ficar sem vigilância”**: Ramachandran, p. 267.
- 227 **A subatividade nessa área**: Carter, pp. 100-103.
- 230 **Dez por cento delas sofrem**: Baker, pp. 46-53.
- 230 **“A Depressão 1.0 podia ser tratada com psicoterapia”**: Ibid., p. 3.
- 232 **De 1% a 3% dos pacientes tratados com ECP**: Carter, p. 98.
- 233 **“as descobertas sobre os canais de cálcio”**: *New York Times*, 26 de fevereiro de 2013, <http://www.nytimes.com/2013/03/01/health/study-finds-genetic-risk-factors-shared-by-5-psychiatric-disorders.html>.
- 233 **“O que identificamos aqui”**: Ibid.

10. A MENTE ARTIFICIAL E A CONSCIÊNCIA DE SILÍCIO

- 239 “as máquinas serão capazes”: Crevier, p. 109.
- 239 “dentro de uma geração... a dificuldade”: Ibid.
- 239 “É como se um grupo de pessoas”: Kaku, p. 79.
- 240 “Eu pagaria um dinheirão por um robô”: Brockman, p. 2.
- 241 **Entretanto, numa reunião particular**: Entrevista com os criadores do Asimo durante uma visita ao laboratório da Honda em Nagoya, Japão, em abril de 2007 para a série da BBC-TV *Visions of the Future*.
- 243 **maravilhado ao ver que um mosquito**: Entrevista com o dr. Rodney Brooks em abril de 2002 para o programa de rádio em rede nacional *Exploration*.
- 249 **Tive o prazer de visitar**: Visita ao MIT Media Laboratory para a série *Sci Fi Science* do canal de TV Discovery/Science em 13 de abril de 2010.
- 251 “Foi por isso que em 2004 Breazeal decidiu”: Moss, p. 168.
- 251 **Na Waseda University**: Gazzaniga, p. 352.
- 251 **O objetivo é integrar**: Ibid., p. 252.
- 251 **Vou lhes apresentar o Nao**: *Guardian*, 9 de agosto de 2010, <http://www.guardian.co.uk/technology/2010/aug/09/nao-robot-develop-display-emotions>.
- 253 “É difícil prever o futuro”: <http://www.cosmosmagazine.com/news/reverse-engineering-brai>.
- 253 **Neurocientistas como o dr. Antonio Damasio**: Damasio, pp. 108-29.
- 262 “Na matemática, não entendemos”: Kurzweil, p. 248.
- 262 “Não é possível haver um teste objetivo”: Pinker, “The Riddle of Knowing You’re Here”, *Your Brain: A User’s Guide*, inverno de 2011, p. 19.
- 263 **Na Universidade Meiji**: Gazzaniga, p. 352.
- 264 “Até onde sabemos, é o primeiro sistema”: Kurzweil.net, 24 de agosto de 2012, <http://www.kurzweilai.net/robot-learns-self-awareness>. Ver também *Yale Daily News*, 25 de setembro de 2012, <http://yaledailynews.com/blog/2012/09/25/first-self-aware-robot-created>. Ver também *Yale Daily News*, 25 de setembro de 2012, <http://yaledailynews.com/blog/2012/09/25/first-self-aware-robot->

[created.](#)

268 **Quando entrevistei o dr. Hans Moravec:** Entrevista com dr. Hans Moravec em novembro de 1998 para o programa de rádio *Exploration*.

268 “**Livres do lento caminhar**”: Sweeney, p. 316.

270 **Quando entrevistei o dr. Rodney Brooks:** Entrevista com dr. Rodney Brooks em abril de 2002 para o programa de rádio *Exploration*.

270 “**Não gostamos de abrir mão**”: TEDTalks,
http://www.ted.com/talks/lang/en/rodney_brooks_on_robots.html.

271 Na Universidade do Sul da Califórnia: <http://phys.org/news205059692.html>.

11. A ENGENHARIA REVERSA DO CÉREBRO

- 273 Q uase simultaneamente, a Comissão Europeia: <http://actu.epfl.ch/news/the-human-brain-project-wins-top-european-science>.
- 279 “é essencial entender o cérebro humano”:
[http://ted.com/talks/henry markram supercomputing the brain s secre](http://ted.com/talks/henry-markram-supercomputing-the-brain-s-secre)
- 279 “Não há uma única doença neurológica”: Kushner, p. 19.
- 281 “estamos longe de brincar de Deus”: Ibid., p. 2.
- 283 “Daqui a cem anos”: Sally Adey, “Reverse Engineering the Brain”, *IEEE Spectrum*, <http://spectrum.ieee.org/biomedical/ethics/reverse-engineering-the-brain>.
- 284 “Os pesquisadores acreditam”:
[http://www.ted.com/talks/lang/en/sebastian seung.html](http://www.ted.com/talks/lang/en/sebastian-seung.html).
- 284 “No século XVII”: [http://www.ted.com/talks/lang.en/sebastian seung.html](http://www.ted.com/talks/lang.en/sebastian-seung.html).
- 285 “O Atlas Allen do Cérebro Humano oferece”: <http://human.brain-map.org>.
- 287 Segundo o dr. V. S. Ramachandran: TED Talks, janeiro de 2010,
<http://www.ted.com>.

12. O FUTURO: A MENTE ALÉM DA MATÉRIA

- 290 **5,8% afirmaram ter tido uma experiência fora do corpo:** Nelson, p. 137.
- 291 **“Eu me vejo de cima”:** Ibid., p. 140.
- 292 **perda temporária de sangue:** *National Geographic News*, 8 de abril de 2010, <http://news.nationalgeographic.com/news/2010/04/100408-near-death-experiences-blood-carbon-dioxide>; Nelson, p. 126.
- 293 **O dr. Thomas Lempert, neurologista:** Nelson, p. 126.
- 293 **A Força Aérea dos Estados Unidos, por exemplo:** Ibid., p. 128.
- 294 **Certa vez, fizemos uma apresentação:** Dubai, Emirados Árabes Unidos, novembro de 2012. Entrevistado em fevereiro de 2003 para o programa de rádio *Exploration*. Entrevistado em outubro de 2012 para o programa de rádio *Science Fantastic*.
- 294 **Em 2055, mil dólares:** Bloom, p. 191.
- 295 **Por exemplo: Bill Gates:** Sweeney, p. 298.
- 296 **“As pessoas que preveem um futuro muito utópico”:** Carter, p. 298.
- 296 **Ele me disse que o zoológico de San Diego:** Entrevista com o dr. Robert Lanza em setembro de 2009 para o programa de rádio *Exploration*.
- 298 **“Devemos ridicularizar os pesquisadores”:** Sebastian Seung, TEDTalks, http://www.ted.com/talks/lang/en/sebastian_seung.html.
- 299 **Em 2008, a BBC exibiu:** <http://www.bbc.co.uk/sn/tvradio/programmes/horizon/broadband/tx/isolati>
- 303 **será possível fazer a engenharia reversa:** Entrevista com o dr. Moravec em novembro de 1998 para o programa de rádio *Exploration*.
- 305 **Do outro lado estava Eric Drexler:** Ver série de cartas em *Chemical and Engineering News* de 2003 a 2004.
- 306 **“Não estou planejando morrer”:** Garreau, p. 128.

13. A MENTE COMO ENERGIA PURA

307 “**Buracos de minhoca, dimensões extras**”: Sir Martin Rees, *Our Final Hour*, Nova York Perseus Books, 2003, p. 182.

14. A MENTE ALIENÍGENA

320 **Até o momento, mais de mil:** Kepler Web Page, <http://kepler.nasa.gov>.

320 **Em 2013, cientistas da Nasa anunciaram:** Ibid.

322 **como ele conseguia distinguir entre mensagens falsas:** Entrevista com o dr. Wertheimer em junho de 1999 para o programa de rádio *Exploration*.

322 **perguntei a ele sobre o “fator risadinha”:** Entrevista com dr. Seth Shostak em maio de 2012 para o programa de rádio *Science Fantastic*.

323 **declarou publicamente que:** Ibid.

323 **“Lembre-se de que é esse mesmo governo”:** Davies, p. 22.

326 **Os historiadores gregos:** Sagan, p. 221.

326 **São Tomás de Aquino dizia:** Ibid.

326 **inteligência pode nos enganar:** Ibid.

327 **“Se o fato de que as bestas não abstraem”:** Ibid., p. 113.

327 **“No mundo surdo e cego do carrapato”:** Eagleman, p. 77.

334 **precisamos expandir nossos horizontes:** Entrevista com o dr. Paul Davies em abril de 2012 para o programa de rádio *Science Fantastic*.

334 **“Minha conclusão é espantosa”:** Davies, p. 159.

339 **“Embora haja uma probabilidade mínima”:** *Discovery News*, 27 de dezembro de 2011, <http://news.discovery.com/space/seti-to-scour-the-moon-for-alien-tech-111227.htm>.

15. OBSERVAÇÕES FINAIS

- 341 Num artigo publicado na revista *Wired*: *Wired*, abril de 2000, <http://www.wired.com/wired/archive/8.04/joy.html>.
- 342 “várias espécies separadas e desiguais”: Garreau, p. 139.
- 342 “Essa utopia tecnológica apela”: Ibid., p. 180.
- 346 “A ideia de que estamos remexendo”: Ibid., p. 353.
- 346 “As tecnologias – como a pólvora”: Ibid., p. 182.
- 349 “Esse *você* que seus amigos conhecem”: Eagleman, p. 205.
- 349 “Nossa realidade depende”: Ibid., p. 208.
- 350 “Algo tão notável como”: Pinker, p. 132.
- 350 criar um planeta gêmeo da Terra: Entrevista com dr. Stephen Jay Gould em novembro de 1996 para o programa de rádio *Exploration*.
- 350 “O *Homo sapiens* é um galhinho”: Pinker, p. 133.
- 351 “nada dá mais sentido à vida”: Pinker, “The Riddle of Knowing You’re Here”, *Time: Your Brain: A User’s Guide*. Inverno de 2011, p. 19.
- 351 “Q ue assombrosa obra de arte”: Eagleman, p. 224.

APÊNDICE

- 362 **muitos (mas não todos) assassinos patológicos:** Entrevista com o dr. Simon Baron-Cohen em julho de 2005 para o programa de rádio *Science Fantastic*.
- 363 **O dr. Michael Sweeney conclui que “os resultados de Libet”:** Sweeney, p. 150.

LEITURA SUGERIDA

Baker, Sherry. "Helen Mayberg." *Discover Magazine Presents the Brain*. Waukesha, WI: Kalmbach Publishing Co., outono 2012.

Bloom, Floyd. *Best of the Brain from Scientific American: Mind, Matter, and Tomorrow's Brain*. Nova York: Dana Press, 2007.

Boileyn-Fitzgerald, Miriam. *Pictures of the Mind: What the New Neuroscience Tells Us About Who We Are*. Upper Saddle River, N.J.: Pearson Education, 2010.

Brockman, John (org.). *The Mind: Leading Scientists Explore the Brain, Memory, Personality, and Happiness*. Nova York: Harper Perennial, 2011.

Calvin, William H. *A Brief History of the Mind*. Nova York: Oxford University Press, 2004.

Carter, Rita. *Mapping the Mind*. Berkeley: University of California Press, 2010.

Crevier, Daniel. *AI: The Tumultuous History of the Search for Artificial Intelligence*. Nova York: Basic Books, 1993.

Crick, Francis. *The Astonishing Hypothesis: The Science Search for the Soul*. Nova York: Touchstone, 1994.

Damasio, Antonio. *E o cérebro criou o homem*. São Paulo: Companhia das Letras, 2011.

Davies, Paul. *The Eerie Silence: Renewing Our Search for Alien Intelligence*. Nova York: Houghton Mifflin Harcourt, 2010.

Dennet, Daniel C. *Breaking the Spell: Religion as a Natural Phenomenon*. Nova York: Viking, 2006.

_____. *Conscious Explained*. Nova York: Back Bay Books, 1991.

DeSalle, Rob e Ian Tattersall. *The Brain: Big Bangs, Behaviors, and Beliefs*. New Haven, CT: Yale University Press, 2012.

Eagleman, David. *Incognito: The Secret Lives of the Brain*. Nova York: Pantheon Books, 2011.

Fox, Douglas. "The Limits of Intelligence." *Scientific American*, julho de 2011.

Garreau, Joel. *Radical Evolution: The Promise and Peril of Enhancing Our Minds, Our Bodies – and What It Means to Be Human*. Nova York: Random House, 2005.

Gazzaniga, Michael S. *Human: The Science Behind What Makes Us Unique*. Nova York: HarperCollins, 2008.

- Gilbert, Daniel. *Stumbling on Happiness*. Nova York: Alfred A. Knopf, 2006.
- Gladwell, Malcolm. *Fora de série – Outliers*. Rio de Janeiro: Sextante, 2009.
- Gould, Stephen Jay. *A falsa medida do homem*. São Paulo: Martins Fontes, 2014.
- Horstman, Judith. *The Scientific American Brave New Brain*. San Francisco: John Wiley and Sons, 2010.
- Kaku, Michio. *A física do futuro*. Rio de Janeiro: Rocco, 2012.
- Kurzweil, Ray. *How to Create a Mind: The Secret of Human Thought Revealed*. Nova York: Viking Books, 2012.
- Kushner, David. “The Man Who Builds Brains.” *Discover Magazine Presents the Brain*. Waukesha, WI: Kalmbach Publishing Co., Outono 2001.
- Moravec, Hans. *Mind Children: The Future of Robot and Human Intelligence*. Cambridge, MA: Harvard University Press, 1988.
- Moss, Frank. *The Sorcerers and Their Apprentices: How the Digital Magicians of the MIT Media Lab Are Creating the Innovative Technologies That Will Transform Our Lives*. Nova York: Crown Business, 2011.
- Nelson, Kevin. *The Spiritual Doorway in the Brain*. Nova York: Dutton, 2011.
- Nicolelis, Miguel. *Beyond Boundaries: The New Neuroscience of Connecting Brains with Machines – and How It Will Change Our Lives*. Nova York: Henry Holt and Co., 2011.
- Pinker, Steven. *Como a mente funciona*. São Paulo: Companhia das Letras, 1998.
- _____. *Do que é feito o pensamento: a língua como janela para a natureza humana*. São Paulo: Companhia das Letras, 2008.
- _____. “The Riddle of Knowing You’re Here.” In *Your Brain: A User’s Guide*. Nova York: Time Inc. Specials, 2011.
- Piore, Adam. “The Thought Helmet: The U.S. Army Wants to Train Soldiers to Communicate Just by Thinking.” *The Brain, Discover Magazine Special*, primavera 2012.
- Purves, Dale et al. (org.). *Neuroscience*. Sunderland, MA: Sinauer Associates, 2001.
- Ramachandran, V. S. *The Tell-Tale Brain: A Neuroscientist’s Quest for What Makes Us Human*. Nova York: W. W. Norton, 2011.
- Rose, Steven. *The Future of the Brain: The Promise and Perils of Tomorrow’s Neuroscience*. Oxford, UK: Oxford University Press, 2005.
- Sagan, Carl. *The Dragons of Eden: Speculations on the Evolution of Human*

Intelligence. Nova York: Ballantine Books, 1977.

Sweeney, Michael S. *Brain: The Complete Mind: How It Develops, How It Works, and How to Keep It Sharp*. Washington, D.C.: National Geographic, 2009.

Tammet, Daniel. *Born on a Blue Day: Inside the Extraordinary Mind of an Autistic Savant*. Nova York: Free Press, 2006.

Wade, Nicholas (org.). *The Science Times Book of the Brain*. Nova York: New York Times Books, 1998.

CRÉDITOS DAS ILUSTRAÇÕES

Página 34: Jeffrey L. Ward

Página 35: Jeffrey L. Ward

Página 37: Jeffrey L. Ward

Página 39: Jeffrey L. Ward

Página 43 (topo): Foto AP/David Duprey

Página 43 (abaixo): Tom Barrick, Chris Clark/Science Source

Página 46: Jeffrey L. Ward

Página 72: Jeffrey L. Ward

Página 73: Jeffrey L. Ward

Página 76: Jeffrey L. Ward

Página 113: Laboratório do dr. Miguel Nicolelis, Duke University

Página 127: Jeffrey L. Ward

Página 250 (topo): MIT Media Lab, Personal Robots Group

Página 250 (abaixo): MIT Media Lab, Personal Robots Group, Mikey Siegel

Título Original
THE FUTURE OF THE MIND
THE SCIENTIFIC QUEST TO UNDERSTAND, ENHANCE, AND EMPOWER
THE MIND

Copyright © 2014 by Michio Kaku

Direitos desta edição reservados à
EDITORA ROCCO LTDA.
Av. Presidente Wilson, 231 – 8º andar
20030-021 – Rio de Janeiro – RJ
Tel.: (21) 3525-2000 – Fax: (21) 3525-2001
rocco@rocco.com.br
www.rocco.com.br

Revisão técnica
VICTOR RIVELLES

Preparação de originais
VIVIAN MANNHEIMER

Coordenação Digital
LÚCIA REIS

Assistente de Produção Digital
JOANA DE CONTI

Revisão de arquivo ePub
ANDRÉ REIS

Edição Digital: julho, 2015

K19f

Kaku, Michio

O futuro da mente [recurso eletrônico] : a busca científica para entender,
aprimorar e potencializar a mente / Michio Kaku ; tradução Angela Lobo. -
1. ed. - Rio de Janeiro : Rocco Digital, 2015.

recurso digital

Tradução de: The future of the mind: the scientific quest to understand,
enhance, and empower the mind

ISBN 978-85-8122-592-0 (recurso eletrônico)

1. Neurociências. 2. Livros eletrônicos. I. Título.

15-23657 CDD: 612.8

CDU: 612.8

O texto deste livro obedece às normas do Acordo Ortográfico da Língua
Portuguesa.

O AUTOR

MICHIO KAKU é um dos físicos mais populares do momento. Ele é professor de física teórica da City University de Nova Yorke e cofundador da Teoria das Cordas. Kaku escreveu diversos livros de sucesso, incluindo *A física do futuro*, *Física do impossível*, *Hiperespaço* e *Mundos paralelos*, todos editados pela Rocco. Este último virou uma série de TV apresentada pelo próprio autor e exibida no Brasil pelo Fantástico, da TV Globo.